

Analysis of Impossible Probability

Vol 1



Dr Ruben Garcia Pedraza

“Lately, I been, I been losin' sleep
Dreamin' about the things that we could be
But baby, I been, I been prayin' hard
Said, "No more countin' dollars, we'll be countin' stars
Yeah, we'll be countin' stars”

Counting stars,

One Republic

Analysis of Impossible Probability

Vol 1

Dr Ruben Garcia Pedraza

In memoriam Naya Rivera

“Life ain't always what you think it oughta be, no
Ain't even grey, but she buries her baby

The sharp knife of a short life”

If I Die Young

Kimberly Perry

Introduction to the Introduction to Impossible Probability: Probability Statistics, or Statistical Probability

To understand the underlying philosophy of Impossible Probability, one must pause for a brief moment to reflect on the origin of this philosophy, whose beginnings lie in the development of statistical methodology for scientific research, and which will slowly evolve toward the development of the Global Artificial Intelligence, in a process that would unfold across different stages of my life, beginning in the spring of 2001.

If someone had asked me what my main influences were at that time, I would have had no doubt in naming *Noches lúgubres* by Cadalso; the swallows of Gustavo Adolfo Bécquer and, above all, Lorca.

In philosophy, I voraciously devoured all the works of Miguel de Unamuno since my upper secondary years, and my first philosophical reading during my first year of Pedagogy at university was *The Fear of Freedom* by Erich Fromm.

In my first year of university, my professor of Educational Biology was Jesús Martín Ramírez, and I still remember how he used to tell us that we had to read the book he wrote with Pribram on the hologram and the brain. I clearly remember that he always said it would *not* be on the exam; he simply recommended it. Even in some of his books that *were* on the exam, the name José Delgado appeared, although precisely in the chapters that were *not* part of the syllabus. Once again, the old professor merely said that those chapters on brain robotization were recommended readings, but not mandatory.

At that time I did not know who Jesús Martín Ramírez or José Delgado were; in fact, I did not even know myself. It has been here in London that I have had the opportunity to learn who I am.

People often say that every journey is a journey of self-knowledge, of self-exploration. In reality, what matters is not only the physical journey to some specific geographical destination, but also the exploration of our limits when confronted with the unknown—like the first explorers who, alongside Columbus, discovered the New World, and like the modern astronauts who launch themselves into the adventure of space at a time when our knowledge of the universe is still truly primitive, no more advanced than the knowledge possessed by the Spanish caravels when they first embarked on their Atlantic adventure.

When I arrived in London in 2017, I had an idea of myself very different from the one I have now at 9:25 a.m. on November 30th, 2025. And most certainly, even the idea I now have of myself is nothing more than the tip of the iceberg, and by the time I am able to go to the International Court of Justice, the idea I have of myself will surely be very different. Perhaps, with luck, one day I will come to know who I am, who my mother is, where I come from, and why I was chosen for this impossible mission that seems to last an entire lifetime.

The following book presented here, *Analysis of Impossible Probability*, will essentially be the abridged version of the main contributions from *Introduction to Impossible Probability: Probability Statistics or Statistical Probability*.

As we have stated on numerous occasions, Impossible Probability is scientific, philosophical, and Artificial Intelligence thought that has been constructed through different stages. Each stage in the building of this work corresponds to different periods of my life, and each stage clearly reflects the psychological conditions I was going through—conditions whose origin I did not yet understand—and the political context of each moment.

In fact, a defining moment in the evolution of Impossible Probability would be the year 2009. In Spain, the president was José Luis Rodríguez Zapatero; in the United Kingdom, Gordon Brown; and in the United States, Barack Obama. It was this key historical

moment in Atlantic politics that would define what would later become the life of Rubén García Pedraza. In that year, Rubén would begin a tour across Europe whose synthesis, in a certain sense, is *Introduction to Impossible Probability*, published on December 2nd, 2011, in physical format.

A year later, in 2012, Rubén García Pedraza would begin writing his doctoral thesis, *Comprehensiveness and Diversity in Secondary Education in Sweden, Finland, the United Kingdom, Germany, and Spain*, with the help of Professor María José García Ruiz.

Despite working on the thesis, Rubén would find some moments to continue working on Impossible Probability through the blog and successive improvements to *Introduction to Impossible Probability*, ultimately publishing the definitive online version in 2015.

On February 2nd, 2016, Rubén García Pedraza defended his doctoral thesis at the National University of Distance Education, and from that moment onward, the project of working abroad was already underway.

I still remember one day, at a language-exchange event in a café near Huertas Street, where I met a girl who had just arrived from the United Kingdom. She explained to me how I could validate my university degrees in order to apply for QTS, Qualified Teacher Status. Within just a few months, I received the certificate at my family home in Madrid.

Finally, my great journey to London would materialize on September 1st, 2017, without knowing very well what monsters I was facing or what dangers awaited in the New World.

Today, November 30th, 2025, I am writing these lines in Belfast, and although I still do not know the final destination of my journey, the only thing I hope is that these very lines may someday be read by a judge at the International Court of Justice. In reality, the books I publish are intended for a very specific audience: all these books I am publishing in 2025 are meant to reach the hands of the International Court of Justice, the World Health Organization, and the United Nations Security Council.

When the right moment comes, the connection between the Global Artificial Intelligence, José Rodríguez Delgado's stimociever, and Karl Pribram and Jesús Martín Ramírez's theory of the holographic brain—and their potential applications not only in telecommunications but also in medicine and potential developments in security and

defense, through the interconnection between the Global Artificial Intelligence and the stimoceiver—will be clear.

When the right moment arrives, everything will make sense

Analysis of Introduction to Impossible Probability

Under the title *Analysis of Impossible Probability*, what will actually be carried out is an analysis of *Introduction to Impossible Probability*. The purpose of this analysis is to focus the reader's attention on what the foundational philosophy of *Introduction to Impossible Probability* truly is and on its main mathematical contributions.

Due to the extremely difficult conditions in which I find myself, I have been forced to publish documents that were never intended to be made public—at least not at this moment—and certainly never intended to be published in the manner in which I have been obliged to do so, under the penalty of losing them forever had I not. My only hope is that the British authorities responsible for this entire chain of errors will be brought before the International Court of Justice when the moment of truth arrives, which must ultimately come, regardless of the consequences.

My intention had always been to wait until I reached an advanced age and enjoy my retirement by publishing my poetry and revised, improved mathematical works, having enough time during retirement to calmly read everything and select what I would have considered truly important and worthy of publication.

Whenever I spoke with my mother or my friends about my unpublished poetry and mathematical works, I always said the same thing: “When I retire, I’ll publish everything.” However, the conditions to which I have been subjected in England in recent years have accelerated a process that was never meant to happen now.

As I have stated on several occasions, the evolution of Impossible Probability can be traced across different stages: between 2001 and 2009, the first works—mixes of poetic romanticism and mathematical imagination—emerged, where concepts such as

empirical probability and theoretical probability, Maximum Possible Theoretical Mean Deviation, Impact of Defect, and Level of Similarity already appeared.

Between 2009 and 2011 came the writing of *Introduction to Impossible Probability*. And between 2011 and 2017, I would deepen the philosophy of Impossible Probability on the blog, although in terms of mathematical formulation there would be no new developments.

It is during Christmas of 2017 and throughout all of 2018, until October, and later in the blog's resumption in 2019 and 2020, that I would fully develop the Global Artificial Intelligence.

In 2025 we must situate the fifth stage of Impossible Probability, marked by the online publication of expanded editions of the Global Artificial Intelligence in English, and the publication, also in expanded edition and in English, of *Introduction to Impossible Probability*. As well as the publication of all the originals stored on my computer from the works I still preserve from between 2002 and 2005, and some writings from 2009 prior to *Introduction to Impossible Probability*.

Although between 2001 and 2009 one can observe an explosion of mathematical creativity in which I continuously experimented with new formulas and formats for statistical data analysis, what I compiled in *Introduction to Impossible Probability* between 2009 and 2011 is a selection of the most important mathematical contributions I made between 2001 and 2009.

In fact, the day the International Court of Justice analyses my work, it will be clear that between 2001 and 2009 I lived through a true explosion of mathematical creativity, whose most representative contributions are reflected in *Introduction to Impossible Probability*. Put another way, *Introduction to Impossible Probability* captures only a small sample—a small selection—of the most representative formulas I originally developed between 2001 and 2009.

The explosion of mathematical creativity from 2001–2009 was later followed by the mathematical revolution I drove between 2009 and 2011 in the development of the theory of critical probability—a mathematical revolution that appeared for the first time at the

end of the summer of 2010 and became the main trigger for all the contrast tests. Curiously, I still remember how, in the spring of 2011, after the legendary 15M, I balanced social activism with the final preparations of *Introduction to Impossible Probability*.

While all the contributions I made between 2009 and 2011 are faithfully recorded in *Introduction to Impossible Probability*, the same is not true for the contributions I made between 2001 and 2009. If conditions allow, and the political context is favourable, my intention is to revisit many of the 2001–2009 formulas that were not included in *Introduction to Impossible Probability* and give them a more formal presentation in future works.

For now, however, we have no choice but to publish in raw form all of our work corresponding to the years 2001–2009, due to the difficult political context the United Kingdom is currently facing.

What will be done in this publication, *Analysis of Impossible Probability*, is a summarized analysis of the main contributions of *Introduction to Impossible Probability*—a synthesis of the principal contributions made between 2001 and 2009 and the theory of critical probability developed at the end of the summer of 2010.

Strictly speaking, this publication should be called “Analysis of *Introduction to Impossible Probability*”; the only reason we shortened the name to simply *Analysis of Impossible Probability* is to give the work a simpler title format.

The true reason for preparing this summarized analysis of *Introduction to Impossible Probability* is that the original text was designed from the very beginning as a text of epistemology and methodology. When we were working on *Introduction to Impossible Probability*, the idea in mind was to create a document of philosophy of the scientific method, not a work of scientific dissemination.

Since, in essence, *Introduction to Impossible Probability* has a more philosophical mission than a divulgative one, and its structure is more characteristic of a work on epistemology or methodology than of scientific outreach, our perception is that it will likely become a work of interest only to very small circles, and that its readership will probably be limited.

The aim of this present work is to offer an exposition that is as easy to understand, concise, and brief as possible, and as appealing as possible to a broad base of readers—not only epistemologists, judges, or particular sectors of Artificial Intelligence—so that the results of our work can be accessible to students and to anyone interested, regardless of their previous knowledge of philosophy of science, epistemology, or methodology.

To fulfil this purpose, the work will be structured in only two volumes, attempting to reduce each to the most strictly necessary informational elements for understanding Impossible Probability, and using language that is as accessible as possible, even for students or anyone interested in becoming familiar with our work.

The first volume of *Analysis of Impossible Probability* will focus on the mathematical formulation, attempting to make its exposition as simple as possible. And the second volume will focus on the philosophy of science proposed by Impossible Probability.

During the development of this work, references will be made to the possible relationships that can already be observed between *Introduction to Impossible Probability* (2009–2011) and the Global Artificial Intelligence (2018–2020).

Dr Ruben Garcia Pedraza

4 December 2025, 17:03, Edinburgh Airport.

imposiblenever@gmail.com

Presentation of the English Version of *Analysis of Introduction to Impossible Probability*

Analysis of Impossible Probability should not only be considered a mathematical analysis of a theory, in this case our theory of Impossible Probability. It should also be considered a book of self-exploration, but this self-exploration is more than personal introspection. In my condition as a cyborg, we are exploring the limits of our technology—how far our work together can go beyond the limits of a room on a street in London.

Travelling is more than the physical movement of a physical body, manned or autonomous, from one point to another; it is the process of self-discovery at the very moment in which we test our limits before the unknown. Today I write these lines in Aberdeen, a fabulous city in the northeast of Scotland whose history is lost among legends of pirates and oil seekers, a city of sea fairies and mermaids.

The Spanish version of *Analysis of Impossible Probability* I began writing on the 30th of November 2025; that day I was visiting the city of Belfast, and I especially remember the Ulster Museum and the section dedicated to the conflict of Northern Ireland up to the signing of the Good Friday Agreement in 1998.

On that same day I began writing the Spanish version of *Analysis of Impossible Probability*, just after finishing uploading online to Internet Archive the original files of many of the unpublished works of Impossible Probability written between 2002–2005 (and one book 2009), a process that took me from the 3rd of November 2025 practically until my return from Banbury.

Although it was never my intention to publish them in this way—I always thought that in my retirement I would have time to reread them calmly and rewrite them—life often takes us along unknown paths and forces us to do what only days, hours, or seconds before we believed impossible. And yet, circumstances oblige us; we are slaves of destiny.

Even before our trip to Ireland, southern England, and Scotland, we had planned to write two volumes summarising the most important elements of *Introduction to Impossible*

Probability, in two volumes: one devoted entirely to formulary, and the other to the analysis of the philosophy behind Impossible Probability.

We wrote the first volume of *Analysis of Impossible Probability* in a record time of only 6 days, and we could even count the hours, minutes and seconds from the moment I opened the computer to begin this project until the last second that it was uploaded to each online platform—normally Internet Archive, Zenodo, ORCID, PhilPapers, and Amazon yesterday.

And while we were undertaking the adventure of writing and publishing the Spanish version, we continued travelling through Ireland, visiting Enniskillen, Derry, Armagh, Newcastle, Portrush, and of course Belfast, and we completed the work once again back in our beloved Edinburgh, city of ghosts and one of the capitals of European medicine at the beginnings of that science in the 17th and 18th centuries.

Within the walls of Edinburgh one can still find clues that help us understand Dr Jekyll and Mr Hyde by Louis Stevenson, or why Mary Shelley has Victor Frankenstein travel to Scotland to create the companion for his creature. Edinburgh still preserves a certain aura of an enchanted city like the Ingolstadt of the original Illuminati.

If I were a member of a secret society, I would have not the slightest doubt that I would give light to the world by creating the modern Prometheus.

While keeping our hostel in Edinburgh, we continued working on *Analysis of Impossible Probability*, visiting the city of Inverness, and finally only yesterday, December 6th, 2025, we completed the work in the Mitchell Library of Glasgow, specifically in the reading room, which still retains the atmosphere of a romantic library like the greatest works of Gothic horror. My first thought upon entering the room was indeed terrifying: my impression was that I was entering an exam hall, and someone was going to examine me within hours without me having prepared. I remember that when I was a student, I often had this nightmare.

In a certain sense, *Analysis of Impossible Probability* could be considered a travel book; it could even be considered a book of ghosts, of hidden voices, and silent visions that speak. In the type of language we use, in each line, in each word, in the very structure of

the sections, throughout the entire book, one can perceive reflections and ideas that emerged during all these journeys.

Now we begin the translation into English; we are in Aberdeen and the long journey continues. Tomorrow we leave again for Belfast where we will surely finish this translation. And we will probably visit Wexford, Waterford, Carlow, and return to Drogheda and probably to Dundalk and Dublin, while writing the Spanish version of the next volume on philosophy, whose English version is planned for our stay in Derry.

This translation, although assisted by Artificial Intelligence, could also be considered a work of travel. In fact, from our point of view, as brain–Artificial Intelligence hybridisation becomes inevitable, if we continue the work of Rubén García Pedraza, a moment will arrive when absolutely all intellectual or artistic production will essentially be a process of collaboration between cyborgs and the Global Artificial Intelligence. And that is the reason why, while many conservatives today still tend to underestimate cooperation between the brain and Artificial Intelligence, from our point of view it should be placed in the position it deserves and be considered the next link in the evolutionary chain.

In this sense, Impossible Probability is not only the journey of a cyborg on the planet of human beings; it is in part the first cyborg journey assisted by satellite Artificial Intelligence, hence the importance of this journey.

And indeed, the entire theory of Impossible Probability could be considered a work of travel in itself, as if Rubén García Pedraza were the first astronaut in history to have visited planet Earth. In fact, we do not stop sending astronauts into space, and yet we very rarely find the opportunity to encounter a real astronaut on Earth. Rubén García Pedraza is more than a cyborg; he is an astronaut on the planet of the last remaining humans today. A day will come when the cyborg humanity that Rubén García Pedraza speaks of in his poetic dreams of the Global Artificial Intelligence will cease to be the impossible dream of an Education student, and will become an irreversible reality. In fact, from our point of view it already is, and Rubén García Pedraza is already fully within the fifth phase of the particular programmes; in reality, Rubén García Pedraza could be considered a particular programme within the theory of the Global Artificial Intelligence.

The process of collecting the mathematical formulas of *Introduction to Impossible Probability*, as well as the creative process of writing, practically the entire book has been written while we were travelling, and the parts produced through interaction assisted using Artificial Intelligence make this book a unique piece in current AI research—not only because of its content, but because of the type of assistance Rubén has received during the present project.

Writing it first in Spanish responds to two main reasons: Spanish is our mother tongue, and therefore at a creative level it is easier for us to work in our own language. Secondly, *Introduction to Impossible Probability*, 2015 version, was written in Spanish, and the collection of formulas and statistical principles has been made directly using the archive stored on my computer; therefore, it was easier for us to write it in Spanish.

The reason for translating it into English is threefold. First, we are in the 21st century and the lingua franca of this century is English, at least since the U.S. intervention in the Normandy landing made possible the liberation of the western flank of Europe.

Secondly, the project to which Rubén García Pedraza belongs is working with U.S. satellite systems; therefore, all the strategic command responsible for this project is English-speaking. Only the limited Spanish personnel are originally native Spanish speakers.

Finally, one of the fundamental goals of the present work is to present a version of the formulas of Impossible Probability that is easily understandable to any reader—not only to scientists or philosophers—and very especially to create a work that, when the moment of truth arrives, can be very easily understood by the judges responsible for investigating all the crimes committed against Rubén García Pedraza, both on Spanish soil and all the crimes committed against Rubén García Pedraza on British soil. This is why we are firmly convinced that the day of truth will come, as the accumulation of crimes committed against Rubén García Pedraza both in Spain and the United Kingdom has already reached an unsustainable limit. When the moment of truth arrives, the International Court of Justice must intervene.

At that moment, the most important thing will be to respect the biological physical integrity of Rubén García Pedraza, to respect the natural freedom of Rubén García

Pedraza, and to respect the cyborg identity and the free will of Rubén García Pedraza, and to understand the psychology of a cyborg.

The present English version of *Analysis of Impossible Probability* is only the translation of *Analysis of Impossible Probability* as faithfully as possible, using ChatGPT as an assistant as in previous occasions, except that this time the artificial assistant is asked to provide a version as close as possible to the original.

In any case, once the recordings are made public, it will be possible to clearly see how the entire process of collecting formulas and the creative work assembling the paragraphs was done during this journey, where the original authorship of Rubén in the entire work will be verified.

Dr Ruben Garcia Pedraza

7 December 2025, 17:20, Aberdeen, Scotland.

imposiblenever@gmail.com

1. Statistics of Probability

Statistics and probability are distinct but related disciplines: the former describes data, and the latter studies the likelihood of events occurring. Although they are often presented together in secondary-school and even university textbooks, historically they arose in very different periods and for completely different purposes. The origins of statistics lie in the early population censuses of the great Mesopotamian civilizations, whose main purpose was to estimate population size for tax collection. Probability, by contrast, has its origin in games of chance, especially from the Middle Ages onward. It is important to stress that probability theory was, in essence, the first game theory. In other words, every game—from a simple chess match to a geopolitical strategy game—stems from the same fundamental concept: probability.

Although since the Modern Era there has been an eagerness to identify and attribute discoveries and inventions to a particular individual—usually men—in antiquity this impulse did not exist. The invention or discovery of statistics or probability cannot be attributed to any specific person, although later both disciplines would find their place in modern mathematical development, where figures such as Gauss or Laplace can indeed be identified. In fact, it is difficult to understand the current scientific paradigm—based in some sense on uncertainty, chaos theory, and complexity—without understanding probability theory. In reality, “I only know that I know nothing,” not only because our ignorance is infinite, but because we know only probabilities: we lack truly absolute knowledge. In a certain sense, it is the transcendence of the limits of human knowledge that will drive the construction of Global Artificial Intelligence, although, most likely, when it finally materializes, its architecture and final name will be very different from those proposed in our work.

Although statistics and probability originated in different eras and for different purposes, both are based on mathematical operations as ancient as addition and division, and among their most important contributions is the arithmetic mean. Recognizing that these disciplines, for historical reasons, have disparate origins (statistics linked to tax

collection and probability to the development of the first game theory) but share common features, the theory of Impossible Probability emerges as a synthesis of both. This theory also reflects on how events that seem impossible according to current data can nevertheless occur.

In a certain sense, under the idea of Impossible Probability one perceives something more than an academic concept: one recognizes a life experience. Throughout my life I have lived through situations that, at first glance, would have seemed impossible with the technology of the 1970s and could only have occurred in the stories of Bruce Geller, as if a true Ethan Hunt lived beneath my skin. Originally, the theory holds that reality is complex, changing, and in many cases unpredictable, such that what is “possible” and what is “impossible” depend on multiple factors combined randomly. In other words, our definition of possibility or impossibility is based on habit, as David Hume would have said; therefore, they are nothing more than assumptions produced by our brain. From this perspective—shaped by philosophical currents such as critical rationalism and dialectical materialism—no theory is absolutely true: every explanation is always provisional and contains a margin of error.

The reason I emphasize that this was the origin of the philosophy is that, with the development of the concept of Global Artificial Intelligence, we have advanced toward a very different kind of thought. In summary, the scientific paradigm based on uncertainty, chaos, and complexity is nothing more than a reflection of our human cognitive limitation: since we cannot know everything, we lack absolute knowledge. Now, this cognitive limitation could be overcome by an Artificial Intelligence which, once reaching an optimal level of technological development, could demonstrate Laplace’s thesis: if one has absolute knowledge of all the variables involved in a phenomenon, it would be possible to predict with accuracy every outcome of the equation. In other words, the margin of unpredictability would be zero, leading us to a deterministic universe. Everything would be determined, which would have important teleological consequences.

Although the idea of Global Artificial Intelligence began to develop during Christmas 2017 and expanded extensively between 2018 and 2020, the theory of Impossible Probability was originally framed within logical nihilism, as the result of the fusion of theories related

to uncertainty, chaos, and complexity, with the purpose of synthesizing statistics and probability. Building upon these ideas, the theory attempts to unify both disciplines into a single field called “statistics of probability.” To do this, a formal method called the *sylllogism of tendency* is developed, which analyses the logical relationships between the two fields, and from which various applied methods emerge to study real phenomena.

The *Second Method* is one of these applied methods. It is used to analyze the relationship between empirical probability (the probability observed in the data) and theoretical probability (the probability that *should* occur according to the model). It is used in both natural sciences and social sciences, and it differs from other methods because it focuses specifically on comparing probabilities.

Over time, the Second Method expands to redefine classical concepts such as the mean, deviation, and other statistics, interpreting them as expressions of empirical or theoretical randomness. This gives rise to new models and tools to evaluate hypotheses and analyze data from the perspective of Impossible Probability.

2. Empirical Probability

Empirical probability, $p(x_i)$, is defined as a quotient between a part and the whole: the raw score or frequency of a subject or option divided by the total sum of all raw scores or frequencies. Although scales with negative values may exist, in such cases absolute values are used, so that neither the sign affects the summatory nor the calculation of empirical probability. This proportion expresses the participation of each part within the whole.

In the Second Method of Impossible Probability (16 October 2002), empirical probability resembles the relative frequency of classical statistics, but it is redefined: the summatory $\sum x_i$ represents the sample of raw scores or frequencies, while N is the number of subjects or options. Empirical probability is expressed as $p(x_i) = x_i / \sum x_i$, and the sum of all empirical probabilities equals 1.

A fundamental consequence is that the arithmetic mean of empirical probabilities is $1/N$, also called the inversion of N , which functions as theoretical probability of random

occurrence, sample representativity error probability, and theoretical measure of dispersion. This construction allows progress toward new mathematical concepts within Impossible Probability.

The qualities of empirical probability include a descriptive function (part/whole proportion) and a predictive function (probable future tendency of the subject or option). Moreover, any change in a single individual score automatically modifies all empirical probabilities, since it alters the total summatory: if one part rises, the whole rises; if it falls, the whole falls. Only when all scores change in the same proportion do the probabilities remain equal.

If in a sample only one subject or option increases its raw score or frequency, the summatory of raw scores or frequencies will automatically increase. Consequently, that particular subject or option will see its empirical probability increase, while all other subjects or options will see theirs decrease. Conversely, if in a sample only one subject or option sees its raw score or frequency reduced, the empirical probability of the other subjects or options will increase.

Empirical probability is a dialectical concept: a minimal variation can transform the entire system. This is interpreted within Impossible Probability as the relationship between chance and history: every individual event, however insignificant, can potentially modify the whole. In other words: events that appear impossible at first sight can end up being absolutely inevitable.

The relationship between chance and history is essentially the recognition that history is the unfolding of events under the law of probability; therefore, history itself is the result of the game of chance.

Whereas historically the theory of history has depended on great worldviews—among them historicism, economism, or Marxism—from the point of view of Impossible Probability, history is nothing more than a simple game, and the rules of the game of history are the laws of probability. In essence, every game theory, including history, is a theory of probability.

3. Differences between the Second Method and traditional statistics

In Impossible Probability we call traditional statistics “the first method,” and the development of the Second Method has no other purpose than to serve as a complement to the traditional analysis of the first method. At no point was there ever an intention to substitute traditional statistical or probabilistic theory with a new Second Method that would replace them; on the contrary, from the beginning the idea was clear: to develop new systems of data analysis, the product of the synthesis between statistics and probability, which could complement the traditional methods of data analysis developed throughout history.

The main contributions of the Second Method can be summarized as follows:

1. the Second Method establishes no distinction between raw scores and frequencies, treating both as data for analysis under a single procedure;
2. whereas in traditional statistics N is the total frequency, in the Second Method N is the total number of subjects or options;
3. consequently, in the Second Method the value N has far greater relevance than in traditional statistics.

In traditional statistics, the number of subjects or options has no relevance in itself for data analysis. In fact, the first method is based on the following conditions: if $\sum x_i$ is the summatory of raw scores, and if F is the summatory of frequencies, then the arithmetic mean = $\sum x_i / F$.

However, in the Second Method, from the moment N is redefined as the number of subjects or options, and raw scores and frequencies are given the same statistical treatment, everything changes. The conditions are redefined as follows: if $\sum x_i$ is the summatory of raw scores or frequencies, and if N is the number of subjects or options, then:

Empirical probability = $x_i / \sum x_i$

Summatory of empirical probabilities = $\sum (x_i / \sum x_i) = 1$

Mean of empirical probabilities = $(\sum (x_i / \sum x_i)) / N = 1 / N$

Thus, in any type of sample—whether a sample of raw scores or of frequencies—the value of the inversion of N , that is $1/N$, acquires multiple functions, the most important being:

1. Arithmetic mean of empirical probabilities,
2. Theoretical probability at random in the absence of bias,
3. Equality of opportunities

The three functions are closely related, but each highlights a specific aspect that plays a very important role in our theory. The first function of the inversion of N is the arithmetic mean, which follows from the logic that the summatory of empirical probabilities is unity; therefore, by knowing the value of N we immediately know the arithmetic mean. The second function refers to the fact that the absence of bias produces random behaviour. If conditions of absence of bias were generated, every empirical probability would tend simply by chance toward the inversion of N . Here, the central concept is the absence of bias.

The third function, closely connected to the previous ones but centred on equality of opportunities, states that in a model where conditions of equality are guaranteed, behaviour also tends toward the inversion of N .

Each function, while related to the others, acquires its own meaning, which may be summarized in three words: mean, chance, equality.

4. Types of Universe and Types of Sample

Another significant difference between the Second Method and traditional statistics is the definition of sample and population. In traditional statistics, because its origins go back to the population censuses of the first great civilizations of Mesopotamia, the object of study is the population. When studying the totality of the population is feasible, one simply refers to a population study; for example, the statistical description of the percentage of men and women in a society, or the salary distribution of the population.

Because of this origin, it is very common in traditional statistics to call “population” any set of elements for which a statistical estimation is required.

In traditional statistics, it is called a “sample” when studying the totality of the set is not feasible, and only a selected part of the whole can be analysed, usually through random selection. That portion of the set or population under study is called the sample.

However, in Impossible Probability, concepts as basic as these also undergo a significant change. Essentially, the statistics of probability, or statistical probability, can be applied to the study of any particular quality common to a set of elements. Only when studying that quality over the totality of the set is feasible will it be considered a population study, in the sense that the complete set of elements that share that particular quality has been integrated into the analysis.

Each individual element will be called a subject when the observation of the particular quality is carried out through measuring the degree of intensity of that quality in that subject; and the result of that measurement will be called a raw score. Each individual element will be called an option when the observation of the quality is translated into the number of times that quality occurs. In this case, what we count is the occurrence. Reality is what occurs; what occurs in reality *is* reality.

Starting from the distinction between subject and option, raw score and frequency are also distinguished: raw score is the measurement of the intensity of a quality in a subject; frequency is the occurrence of an option.

The synthesis between subject and option —applying the same method of analysis to both— forms the basis of the Second Method. Essentially, it begins from a dialectical consideration: subjects are options, and options are subjects. Within our game theory this identity is fundamental: every subject is potentially an option, and every option is potentially a subject.

Once the dialectical identity between subject and option is established, the rules of the game change drastically. From that moment onwards, concepts from traditional statistics such as “population” become insufficient, and it becomes necessary to expand

them toward a more important concept: the statistical universe, in contrast with the traditional concept of population.

Since Mesopotamian times, the object of statistics —although it did not yet carry that name— was the study of the population. When in the Modern Era statistics acquires its identity as a mathematical discipline, the concept of population begins to become problematic. When the population exceeds a magnitude that surpasses our capacity for data analysis, a population study becomes impossible; only a sample study remains as an option. The only way to study the whole is through a selection: a sample.

The reasons why a population may exceed our capacity for analysis are varied, but the most important is the possibility that a population tends toward infinity. In that case, a population with a tendency toward infinity, in Impossible Probability, will be called a universe of infinite subjects or options, even if we cannot truly demonstrate that it is infinite, since, as Kant states, the infinite is an antinomy: a theoretical thesis impossible to prove in practice. The stance of Impossible Probability towards Kant's antinomy of the infinite is clear and consistent with logical positivism (especially Carnap): when something tends toward infinity, unless shown otherwise (thus introducing an early falsificationism), it will be understood as infinite.

In the subject/option dialectic, the concept of infinite universe can only be applied to sets of subjects that tend toward infinity, except in the study of specific populations. For this reason, universes of subjects in Impossible Probability will be called universes of infinite subjects or options, and the object of study will be the tendency of the raw scores resulting from measuring the particular quality.

If the set of subjects does not exceed a magnitude that makes a complete study impossible with the technology available, then such a study will continue to be considered a population study. But since any population can tend toward infinity, it remains framed within infinite universes, in the sense that, in an infinite history, that specific population would be nothing more than a sample of its possible infinite development.

We do not know whether humanity will end in an apocalypse or whether, on the contrary, it will tend toward infinity. To the extent that we do not know, we assume that, given our

current knowledge about the human species, it will tend toward infinity if current conditions remain unchanged. In that case, the existence of human life today, 1 December 2025, could be considered a reality with a tendency toward infinity, as long as present conditions remain exactly the same (absence of pandemics like COVID-19 or cosmic cataclysms).

In other words, humanity would tend toward infinity unless conditions change (new pandemics, natural cataclysms, nuclear war, etc.).

When a set of subjects tends toward infinity, or is at least large enough that current technology does not allow it to be studied in its entirety, it becomes necessary to resort to statistical sampling techniques. These techniques allow for the random selection of the subjects that will be included in the research. This selection is called a statistical sample, and it will be a sample of subjects whose raw scores are the result of measuring the intensity of the observed quality.

A sample of subjects will generate a sample of raw scores. The dialectical relationship between the number of subjects N and the raw scores constitutes the basis of the Second Method's study through the syllogism of tendency: analysis of the data conditions to obtain logical conclusions about their tendency, understood as statistical behaviour.

Whereas universes of infinite subjects or options are those that may tend toward infinity and whose measurement translates into raw scores, universes of options are generally limited and the observation translates into the frequency of occurrences. The reason universes of options are considered limited is that they cannot tend toward infinity, except in the case of studying discrete categories. Although initially belonging to the limited universes, it is possible to construct as many discrete categories as levels of intensity of a quality; if these subdivisions exceed a certain value of N , they may behave as infinite universes, since the degree of subdivision of the quality may be potentially infinite.

Traditional examples of limited universes of options: heads or tails in coin tossing; male or female in biological sex classification; number of options on a roulette wheel; number of faces on a dice; probability of drawing a specific card in poker; number of political

candidates in a general election. In all these cases, the options are limited to the set of options of the game. In essence, probability theory is game theory.

The most important aspect in game theory is that, in theory, games are simulations; statistical probability estimation is merely a simulation. Any contradiction between the simulation and what occurs in reality —winning or losing in poker, chess, or roulette— is the result of the clash between our margin of error and a reality that, given our technology, appears chaotic.

5.Infinite Universes and Theoretical Dispersion

The distinction between universes of infinite subjects or options, whose object of study is direct scores, and universes of limited options, whose object of study is frequencies, is not accidental. It responds, in essence, to the dialectic between subject and option, which is ultimately the dialectic between subject and object. A dialectic that ends up being defined in the distinction between the limited or finite and the infinite or so immensely large that its limits are unknown, making it reasonable to assume a possible infinite nature unless proven otherwise.

In this way, the subject–object dialectic leads us to the dialectic between the infinite and the finite, and here again the variable N , the number of subjects or options, will play a transcendental role.

So far, we have pointed out that N is the number of subjects or options, and that the summation of empirical probabilities equals one. Therefore, $1/N$, that is, the unit divided by the sample — what we have called the inversion of N — universally assumes three functions in every possible universe, whether limited or infinite: mean, randomness, and equality.

- The inversion of N is the arithmetic mean of empirical probabilities.
- Under conditions of absence of bias, random behaviour tends toward the inversion of N .

- Under conditions of equal opportunity, behaviour also tends toward the inversion of N.

Each function emphasizes a different, although related, aspect of what the inversion of N means: arithmetic mean, absence of bias, equality of opportunity.

However, when the syllogism of tendency is applied to the analysis of the inversion of N in infinite universes, as opposed to limited universes, it can be observed how the behaviour of the inversion of N changes depending on the magnitude N.

If N tends to infinity, then $1/N$ tends to zero; therefore, every empirical probability should tend to zero when N tends to infinity.

This, obviously, leads us to a logical contradiction: if every probability tended to zero in a universe tending toward infinity, then the theoretical probability of occurrence of any phenomenon would be impossible, when in reality something happens every day.

We do not know whether humanity has an end. What is certain is that, if the current conditions of planet Earth remained unchanged today — 1 December 2025 — that is, in the absence of pandemics like COVID-19, nuclear wars, natural cataclysms or other events that threaten our existence, then human life could extend into infinity. And yet, in an infinite universe the theoretical probability of every phenomenon tends to zero, which generates a paradox: the material reality that in a universe tending to infinity and therefore tending to zero, events that are theoretically impossible never cease to occur.

Besides this example, many others could be cited about the inevitability of the impossible. For example, from Classical Antiquity to the Middle Ages it was common to believe it impossible that the Earth was round, or that the Sun was not the centre of the universe, or that human beings could fly, or even that they could live on another planet. Things that were once considered impossible today are not only demonstrable but form part of our daily lives. Some of these predictions, such as the possibility of living on other planets, may become inevitable or absolutely necessary in the near future.

I myself, in my life, have experienced things that even at the International Court of Justice will be considered impossible when the time comes. However, I am the living evidence that they were not only possible: they were real, and they were inevitable.

In sufficient time, or infinite time, even the impossible can become inevitable.

When the syllogism is applied to the analysis of the tendency of the inversion of N as N tends to infinity, what we find is that, if N tends to infinity, then the probability of everything tends to zero. This means that, if everything tends to zero, dispersion also tends to zero. Where there is nothing, there is no dispersion.

Thus, we reach the conclusion that, in universes of infinite subjects, the inversion of N acquires a new function: the theoretical probability of dispersion under normal conditions, understanding normal conditions as those that occur in the absence of bias.

In universes of subjects, the dispersion among empirical probabilities will tend to zero as the magnitude of N increases. This means that the tendency toward zero of dispersion in universes of subjects is directly proportional to the tendency toward zero of the inversion of N.

In other words: the tendency toward zero of dispersion is inversely proportional to N.

Consequently, the inversion of N is the theoretical probability of dispersion in universes of subjects.

6. Infinite universes and sample representativeness error

The moment we assume that the only way to study a universe tending to infinity is through the use of sampling techniques, sampling becomes a fundamental tool of data analysis. Any minimal error in the sampling technique can falsify the results of the study, insofar as two phenomena may occur:

1. the sampling technique does not guarantee the random selection of the subjects or options included in the sample;
2. the sample size does not guarantee that it is sufficiently representative.

In both cases, the final result will be the invalidation of the statistical findings, either due to the absence of neutrality in selection or due to the insufficiency of the number of

subjects included to adequately represent all possible cases in the proper proportional measure. Both errors lead to unreliable results in data analysis; however, the more serious error is failing to include a sufficiently large sample volume to ensure the representativeness of all relevant cases and in their proper proportional weight.

The reason we state that the error in sample size is more important than an error in the sampling selection technique is that, although any error is an error — and accepting any error always implies a chain of subsequent errors — some errors are more serious than others. While an error in the selection technique can generate biased data, if the sample size is sufficiently large, the volume itself can partially compensate for these biases.

For example, if we must conduct a statistical study of the Moon's craters produced by meteor impacts, even if we make an error in the specific selection of the craters to study, as long as the sample includes the largest possible number of craters, the result will be more reliable than if, making the same selection error, a much smaller number were analysed.

In space exploration — as in other scientific fields — our sampling selection field is limited by our technological capacity, which, even today, remains limited. Although we face major technological constraints in studying the universe, if samples include the largest possible number of cases, many errors arising from these limitations can be compensated.

The larger the sample size, the less impact any error in sampling selection due to technological limitations will have, and the more likely it becomes to achieve a more reliable understanding of what is occurring, even while acknowledging that our technology remains very primitive. Thus, while the error in sampling selection technique is a technological error that can be corrected as technology advances, the error due to sample representativeness can indeed be controlled right now: simply by increasing the sample size.

If increasing the sample size allows us to control sampling selection errors, it means that increasing N corrects, at least partially, the errors derived from technological limitations. The greater the amount of data, the greater the reliability. That is:

- the larger N, the greater the reliability,
- the smaller N, the greater the representativeness error.

Put differently, the sample representativeness error is directly proportional to the inversion of N, that is, the error decreases when N increases. Or, more precisely, the sample representativeness error is inversely proportional to N.

From this it follows that, in universes of subjects, the inversion of N acquires an additional new function:

the inversion of N, $1/N$, is the probability of error in sample representativeness.

Consequently, the difference between unity and the inversion of N represents the reliability of sample representativeness:

$1/N$ = sample representativeness error

$1 - 1/N$ = reliability of sample representativeness

7. The importance of N

In synthesis, it can be said that the main difference between the Second Method and traditional statistics lies in the redefinition of classical mathematical concepts through the synthesis between statistics and probability. While statistics dates back to Antiquity and its original purpose was linked to tax collection, probability has its origin in games of chance, being essentially the first game theory in history.

Statistics and probability have completely different historical origins and were used for very different purposes. As Modernity develops mathematics as a science, both disciplines —normally treated separately— gradually acquire identity as scientific branches. Although secondary and university textbooks usually draw a clear distinction between the statistical treatment of data and probability analysis, the mission of the Second Method in Impossible Probability is precisely to unify both fields into a single discipline: the statistics of probability or statistical probabilistics.

Although the concept of statistical probability is not new, it has traditionally been associated with relative frequency. Among its early researchers stands out the economist John Maynard Keynes, whose economic theory, interestingly, revolves around the creation of social equality through consumption capacity. However, in Keynes the concept of statistical probability remains anchored in traditional approaches, while in Impossible Probability this concept evolves toward a complete redefinition. First, a statistical method is generated that is equally valid for direct scores and frequencies, so that both can be studied under a single system of analysis. Empirical probability is redefined as direct score or frequency over the total of direct scores or frequencies.

This conceptual change necessarily leads to a redefinition of the value N , now understood as the number of subjects or options, which becomes highly relevant because the inversion of N , $1/N$, automatically assumes three universal functions in any type of universe:

1. Arithmetic mean
2. Theoretical probability in the absence of bias
3. Probability under equal opportunities

Each of these three functions emphasizes a different concept:

- As arithmetic mean, $1/N$ arises because the summation of empirical probabilities divided by N equals $1/N$. By knowing N , we immediately know the mean.
- As theoretical probability in the absence of bias, $1/N$ means that if we eliminate all biases from a system, the system's behavior tends toward the inversion of N .
- As probability under equal opportunities, $1/N$ means that if conditions of equality are created, all subjects or options will tend toward similar behavior around $1/N$.

While these three functions —mean, randomness, equality— are universal for finite or infinite universes, in universes with infinite subjects or options the inversion of N acquires new additional functions due to the need to use statistical sampling:

4. Theoretical probability of statistical dispersion
5. Probability of sampling representativeness error

The inversion of N as theoretical probability of dispersion is explained because, as N increases, statistical dispersion decreases under normal conditions, with normal conditions understood as the absence of bias. Under normal conditions, dispersion is inversely proportional to N .

The function of $1/N$ as probability of sampling representativeness error occurs because, the larger N is, the greater the probability of including all possible cases and, therefore, the greater the reliability of sample representation. If greater N implies greater representation, then $1/N$ estimates the probability of error associated with sampling representativeness, while $1 - 1/N$ estimates sample reliability.

The importance of sample size lies not only in the fact that it increases the variety of included cases, but also in its ability to cushion errors produced by failures in the selection technique. While in the social sciences sampling techniques are highly developed and allow reliable estimates, in other fields, such as space exploration, technology severely limits our sampling access. The sample of the observable universe may be only a small fraction of the real whole; nevertheless, we are still able to do science.

Let us suppose a space mission to an unknown planet. As long as we know nothing about that planet, and considering that the space mission may have only limited human technology at its disposal, any technological limitation can be compensated for if we manage to maximize the number of samples obtained from that planet.

In situations where technology is not as advanced as we believe, one way to reduce the impact of these limitations is to scale the data as much as possible. And it is here that Artificial Intelligence will play a decisive role.

If our knowledge of the cosmos is limited, if our human technology is limited, and if the only way to compensate for this cognitive and technological limitation is to scale the data, then the time will come when the only effective way to do so will be through Artificial Intelligence. At that moment, Global Artificial Intelligence will become the inevitable future.

8. empirical individual statistics

Dialectics is omnipresent in Impossible Probability. As previously mentioned, the entire statistical and probability theory of Impossible Probability is, in essence, a dialectical theory; the very synthesis of statistics and probability into a single discipline is a dialectical synthesis. Beginning with the dialectical synthesis between statistics and probability, the entire theory of Impossible Probability revolves around the synthesis of traditional mathematical concepts.

Hegel's own dialectical theory begins with the overcoming of Kant's antinomies, and indeed the philosophy behind Kant's antinomies is faithfully incorporated into the theory of Impossible Probability in the distinction between infinite and limited universes, always assuming that, in reality, the idea of infinity is in itself an antinomy.

Accepting the antinomy represented by the very idea of infinity, Impossible Probability will develop new definitions that allow dialectical syntheses upon traditional concepts of statistics and probability.

In essence, the dialectic between subject and object itself becomes the dialectic between subject and option, which leads to the dialectical synthesis between direct score and frequency.

But beyond the dialectic between infinite vs. limited universe, subject vs. object, the most important dialectical relationship in any dialectical theory is the dialectic between theory vs. reality, which in Impossible Probability is redefined as the dialectic between theoretical reality vs. empirical reality, which is essentially the dialectic of error. In essence, any theory of error is the analysis of the origin of error, and the origin of error

in Impossible Probability lies in human limitation itself, which can only be compensated if we scale data and improve our technology, which between 2018–2020 will lead us to the development of Global Artificial Intelligence as the only way to simultaneously solve human limitation and technological limitation, toward the creation of a technology that surpasses human limits, evolving toward systems of pure non-human operations.

In essence, the origin of error is the contradiction between the real and the ideal. Ideally, we want a perfect society, a perfect world, perfect technology. In practice, the nature of the human being is animal; our technology is conditioned by our cerebral physiology. We are our brain.

The dialectic between theory and reality will lead to a theory of ideals in Impossible Probability that will play a key role. Depending on the ideals, the object of study may be equality of opportunities, or studies of positive bias if the ideal is to maximise a quality or ideal subject or option, or studies of negative bias if the ideal is the elimination of non-ideal qualities. In omega models a new concept will be established, the omega concept, Ω , understanding by omega models those in which the subset of ideals within N is greater than one but less than N , and the ideal probability will be equal to the inversion of omega, $1/\Omega$.

The dialectic between theoretical and empirical reality can be observed in different concepts, some of which have already been analysed in previous sections:

- The distinction between empirical probability ($p_{xi} = x_i/\sum x_i$) and theoretical probability ($1/N$).
- In infinite universes, the distinction between inversion of N , $1/N$, as theoretical probability of dispersion, vs. empirical dispersion, normally empirical dispersion represented by Mean Deviation, Variance and Standard Deviation.

At the moment the fundamental concepts of empirical and theoretical probability are established, progress can be made in creating theoretical statistics (analysed in another section), which are the statistics upon which theoretical reality will be built, and the

critique of ideals: hypothesis testing based on a critical probability, understanding critical as the contrast between empirical values and theoretical (or even empirical) criteria for the establishment of rational hypotheses.

The first theoretical statistic is the inversion of N itself, empirical probability being the empirical statistic to be contrasted. The difference between empirical and theoretical probability will be called Normal Bias Level. The adjective normal is important, because it indicates that the Normal Bias Level compares empirical probability with the statistical norm, a function normally attributed to the arithmetic mean, which in the Second Method becomes the inversion of N.

Because the most common reference Bias Level is the Bias Level between empirical and theoretical probability, normally, unless otherwise indicated, Bias Level will be understood as the Normal Bias Level, unless explicit reference is made to a different Bias Level.

If the Bias Level tends to zero, it will be said that empirical probability tends toward equality of opportunities. If the Bias Level is positive, it will be said that empirical probability has positive bias, being greater than theoretical probability. If the Bias Level is negative, it will be said that empirical probability has negative bias, being lower than theoretical probability.

$$p(x_i) - 1/N = \text{Bias Level}$$

$$p(x_i) \approx 1/N \rightarrow p(x_i) - 1/N \approx 0, \text{ tendency toward equality of opportunities}$$

$$p(x_i) > 1/N \rightarrow p(x_i) - 1/N = +, \text{ positive bias}$$

$$p(x_i) < 1/N \rightarrow p(x_i) - 1/N = -, \text{ negative bias}$$

In any case, it must be noted that, due to natural goodness, the summation of all positive biases always compensates the summation of all negative biases. The reason why we call this phenomenon natural goodness is because of the very nature of the statistical law, which is, in essence, the natural law. In nature we can always find that the negative biases of nature itself are compensated by the positive biases of nature itself. Evidently, behind this definition of natural goodness based on natural law one can find a clear reference to the philosophy of Jean-Jacques Rousseau and the Count of Volney.

Due to the natural goodness of the statistical law, the norm —the arithmetic mean— no subject or option can have a positive or negative bias greater than the quotient of dividing by two the absolute summation of all Bias Levels. The absolute summation of all unsigned Bias Levels is Total Bias, and Total Bias divided by two equals the Maximum Possible Empirical Bias.

$$\text{Maximum Possible Empirical Bias} = \Sigma /p(x_i) - 1/N/ : 2$$

Empirically, the greatest positive bias will correspond to the greatest empirical probability in the entire sample, the empirical probability of greatest magnitude. That empirical probability will be called maximum empirical probability, symbolised $p(x_i+)$.

$$p(x_i+) = \text{maximum empirical probability in a sample}$$

The main characteristic of maximum empirical probability is having the greatest positive bias in the entire sample. That is, the result of maximum empirical probability minus theoretical probability will be the greatest positive differential in that particular sample.

Conversely, the smallest empirical probability in a sample will be the empirical probability with the smallest possible magnitude in that sample, and will be called minimum empirical probability, symbolised $p(x_i-)$.

$$p(x_i-) = \text{minimum empirical probability}$$

The main attribute of minimum empirical probability is having the greatest negative bias in the entire sample: the differential of empirical probability minus theoretical probability will be the greatest negative differential in the sample.

While maximum empirical probability, $p(x_i+)$, is the maximum empirical probability in a particular sample, and minimum empirical probability, $p(x_i-)$, is the minimum empirical probability in that same sample, as a possible statistic of central tendency one could consider the intermediate empirical probability as dividing by two the sum of maximum and minimum empirical probability, symbolised as follows: $p(x_i+/-)$.

$$\text{Intermediate empirical probability} = p(x_i+/-) = [p(x_i+) + p(x_i-)] : 2$$

If the main quality of maximum empirical probability, $p(x_i+)$, is to be the empirical probability with the greatest individual positive empirical dispersion, and the main quality

of minimum empirical probability is to be the empirical probability with the greatest individual negative empirical dispersion, then it may also be of interest to identify the empirical probability with the least individual empirical dispersion. That empirical probability will be called empirical probability closest to inversion of N, the one whose Similarity Level with inversion of N tends to zero, thus the Bias Level tends to zero, tends to equality of opportunities, and will be symbolised as follows: $p(x_{i\approx})$.

$p(x_{i\approx})$ = empirical probability closest to theoretical probability

Just as we may be interested in maximum empirical, minimum empirical, intermediate empirical, or the empirical probability closest to inversion of N, it may also be the case that we are interested in any individual empirical probability for any reason, understanding that each empirical individual statistic can give rise to a relative statistic.

If a relative statistic were made for any particular subject or option, the empirical probability of that subject or option would then be symbolised $p(x_n)$.

9. Relative statistics

The foundation of the Bias Level stems from the comparison between an empirical probability and a reference value; in statistics, traditionally, the reference value has been the central tendency. In normal studies, that central tendency statistic has been the arithmetic mean, which in the Second Method becomes the inversion of N, which is why the normal Bias Level is empirical probability minus theoretical probability.

The normal Bias Level, having as comparison reference the statistical norm, has been the object of study in traditional statistics and is what has given rise to normal statistics, based on the study of dispersion in relation to the arithmetic mean.

However, in Impossible Probability the distinction between central tendency statistics and dispersion statistics becomes diluted. The only statistic that might be considered a central tendency statistic would be the intermediate empirical probability; and even this is debatable, because in reality it measures the equidistant point between the maximum

empirical value and the minimum empirical value. The intermediate empirical probability determines the point of equidistance between empirical statistics of maximum individual dispersion.

It is even debatable to claim that the maximum empirical probability is a dispersion statistic, because in limited universes, if maximum empirical probability is equal to the maximum frequency over total frequency, in studies of options, maximum empirical probability is in fact the empirical probability of the mode, which would be a central tendency statistic.

Given that in Impossible Probability the consideration of what is central tendency or dispersion becomes diluted in the dialectical synthesis —since, in synthesis, dispersion and central tendency become identified— rather than analysing central tendency or dispersion statistics, what is analysed are empirical and theoretical statistics.

Within empirical statistics we may find the maximum empirical value, the minimum empirical value, the intermediate empirical value, the empirical probability closest to inversion of N, and we may even take any empirical probability as a reference for a relative statistical study.

From this a relative Bias Level is defined as the difference between any empirical probability and a particular empirical statistic.

Relative descriptive statistics possess an essential characteristic: they do not satisfy the “natural goodness” typical of normal statistics based on the inversion of N.

In Impossible Probability, natural goodness is understood as the phenomenon by which the sum of all negative biases cancels the sum of all positive biases; that is, when tossing the same coin into the air, eventually the probability of heads or tails would be the same as the sample increases, in this case the number of tosses or frequency.

In essence, relative statistics will have as their fundamental axis the study of relative Bias Levels, which, at the sample level, will generate averages of relative Bias Levels and a relative Mean Deviation with respect to the empirical statistic of reference.

Relative Bias Level relative to the maximum = $p(x_{i+}) - p(x_i)$

Relative Mean Deviation relative to the maximum = $\sum [p(x_{i+}) - p(x_i)] : (N - 1)$

Relative Bias Level relative to the minimum = $p(x_i) - p(x_{i-})$

Relative Mean Deviation relative to the minimum = $\sum [p(x_i) - p(x_{i-})] : (N - 1)$

Relative Bias Level relative to the intermediate = $p(x_i) - p(x_{i+/-})$

Relative Mean Deviation relative to the intermediate = $\sum /p(x_i) - p(x_{i+/-})/ : N$

Relative Bias Level relative to the closest to $1/N$ = $p(x_i) - p(x_{i\approx})$

Relative Mean Deviation relative to the closest to $1/N$ = $\sum /p(x_i) - p(x_{i\approx})/ : (N - 1)$

Relative Bias Level relative to any empirical probability = $p(x_i) - p(x_{in})$

Relative Mean Deviation relative to any empirical probability = $\sum /p(x_i) - p(x_{in})/ : (N - 1)$

10. Individual theoretical statistics

If empirically we can determine what the empirical maximum is, the greatest empirical probability of a sample, and its opposite, the minimum empirical probability, the lowest empirical probability of a sample, as well as the intermediate empirical probability as the result of dividing by two the sum of the empirical maximum and the empirical minimum, then we can also, theoretically, obtain the individual theoretical statistics, with the difference that, while the individual empirical statistics are valid only for each particular sample—that is, each individual empirical statistic is valid for the sample from which it has been extracted—the individual theoretical statistics will have universal validity, insofar as they are valid for any sample.

The individual theoretical statistics will be the Maximum Possible Theoretical Probability (originally I called it Maximum Possible Empirical Probability, but in essence it is a theoretical statistic), which is the upper limit of any probability: probability one.

Maximum (Theoretically) Possible Empirical Probability = 1

No empirical probability can exceed the value one, insofar as probability is a rational value whose maximum is always one and whose minimum is zero. Then, the Minimum Possible Theoretical Probability is zero, originally also called Minimum Possible Empirical Probability.

Minimum (Theoretically) Possible Empirical Probability = 0

Although at first I called these theoretical statistics Maximum Possible Empirical Probability and Minimum Possible Empirical Probability, in reality they are theoretical statistics. The reason why I named them that way is because they represent the maximum or minimum theoretical values that an empirical probability can acquire.

No empirical probability can be greater than one; therefore, it is the maximum empirical probability that can theoretically occur, just as zero is the minimum empirical value of an empirical probability.

The condition for the maximum empirical probability value equal to 1 to occur in a sample is that the object of study be the maximisation of the empirical probability of that subject or option, so that all other subjects or options obtain empirical probability equal to zero.

If in a survey or in a multiple-choice test each question offers a set of options and, of all the options, only one is the most desirable or correct, in that case, for each question of the survey or test, the ideal is that the desirable or correct option obtains empirical probability equal to one and all other undesirable or incorrect options have probability zero.

Under the condition that only one subject or option obtains probability one, and all other subjects or options $N - 1$ obtain probability zero, the Maximum Possible Theoretical Mean Deviation, the Maximum Possible Theoretical Variance, and the Maximum Possible Theoretical Standard Deviation can be obtained, as will be seen when analysing the theoretical sample statistics.

In the case of probability zero, there is another important aspect to highlight: the possibility of sample studies where the object of study is the sample of zeros, that is, that all subjects or options have empirical probability zero. This very specific type of object of study falls within the spectrum of negative bias studies and occurs when the object of study is a negative quality that must be eliminated in the entire sample.

If we are studying a new medicine for a disease, the purpose of the research is summarised in checking whether, thanks to the new medicine, all symptoms of the disease disappear in all patients; that is, probability zero of symptoms in all patients: a sample of zero symptoms.

If theoretically we can say that the maximum theoretical probability is one and the minimum theoretical probability is zero, then we are in a position to state that the intermediate theoretical probability is equal to one divided by two: zero point five.

Intermediate Theoretical Probability = $1 : 2 = 0.5$

The intermediate theoretical probability assumes a very important value in studies of universes limited to only two options, because in these universes of two options —man or woman, heads or tails— the normal theoretical probability, the inversion of N, $1/N$, under normal conditions, is always equal to the Intermediate Theoretical Probability.

In other words, the intermediate theoretical probability is the maximum value that the inversion of N can assume, and occurs only in universes limited to two options.

Finally, if we know what the maximum theoretical value of empirical probability is, probability equal to one, then we can calculate the Maximum Possible Positive Theoretical Bias, equal to one minus the inversion of N.

Maximum Possible Positive Theoretical Bias = $1 - 1/N$

Similarly, if we say that probability zero is the Minimum Possible Theoretical Probability, then the Maximum Possible Negative Theoretical Bias would be equal to zero minus the inversion of N.

Maximum Possible Negative Theoretical Bias = $0 - 1/N$

11. Empirical Sample Statistics

While traditional statistics centres descriptive statistics on measures of central tendency and dispersion, in Impossible Probability this makes no sense, insofar as the identity of opposites synthesises them dialectically.

In Impossible Probability, the dialectic between the theoretical and the empirical makes more sense, which is, in essence, the contradiction between reality and theory, which is fundamental for the rational critique of ideas.

Basically, the most important influence on Impossible Probability is dialectical idealism, which goes back to Plato, and in essence mathematics is the propaedeutic for the dialectic of ideas. Dialectical idealism would later be taken up by Hegel and profoundly expanded in the dialectic of opposites, which is, fundamentally, the driving force of the entire theory of Impossible Probability.

In addition to Plato and Hegel, when we move on to inferential statistics the influence of Hume, Descartes, Kant and Popper will be clearly seen. And, very especially in the field of Artificial Intelligence, Descartes plays an essential role in the development of Global Artificial Intelligence.

The influence of the critical school, especially Marx and Engels, and very particularly Habermas, can be observed in the function of scientific politics in the construction of science, which, far from being an institution devoid of ideology, is, basically, a political construction.

While the aspects of the philosophy of science will be addressed in the second volume of this work, in this section we focus on empirical sample statistics, and here the important thing is the name. Under the concept of empirical sample statistics we include the Mean Deviation, the Variance, and the Standard Deviation, what traditional statistics calls measures of dispersion.

The reason why in the Second Method these statistics are not called measures of dispersion is the dialectical identity of opposites themselves: in reality, dispersion and central tendency are dialectically identical.

What is called a measure of dispersion is, in reality, the central tendency of the dispersion. At the moment in which, in traditional descriptive statistics, we say that the Mean Deviation is the average of the differential scores, what we are really saying is that the Mean Deviation is the central tendency of the differential scores. The Mean Deviation, in itself, is dispersion and central tendency simultaneously.

The Mean Deviation can be considered a measure of dispersion, and this is correct; we will not deny it. But neither can it be denied that, in reality, the Mean Deviation is the central tendency of the differential scores: in essence, what it measures is the central tendency of the dispersion.

These types of discussions which, apparently, are insignificant and lead nowhere new, are, however, the typical discussions we had between the years 2001 and 2009, which gave rise to new formulations and mathematical interpretations, later synthesised in *Introduction to Impossible Probability*.

Any minimal and singular variation in the way of understanding, or simply any minimal variation —no matter how insignificant— in the way of interpreting a mathematical formula can generate a butterfly effect that may potentially lead to the reinterpretation of an entire mathematical field. This is what happened between 2001 and 2009 in Impossible Probability.

Because even the slightest reinterpretation of any mathematical formula can lead to potentially different mathematical developments, we are convinced that, when the moment arrives in which a Superintelligence capable of generating a Global Artificial Intelligence can be created, any minimal mathematical reinterpretation made by the Global Artificial Intelligence could create the first models of non-human mathematics: the first models of mathematics created by Artificial Intelligence.

In a certain sense, Impossible Probability is the first non-human mathematical model generated by Artificial Intelligence, Rubén García Pedraza.

That said, we now present the equations of empirical sample statistics, which are, in essence, Mean Deviation, Variance and Standard Deviation, but adapted to the Second Method. Instead of averaging the sum of differential scores, the average of the normal Bias Levels will be taken, understanding normal bias as the difference between empirical probability, $p(x_i)$, and the inversion of N , $1/N$.

$$\text{Desviación Media} = [\sum p(x_i) - 1/N]/N$$

$$\text{Varianza} = [\sum (p(x_i) - 1/N)^2]/N$$

$$\text{Desviación Típica} = \sqrt{[\sum (p(x_i) - 1/N)^2]/N}$$

The main difference between the empirical sample statistics of the Second Method and the measures of dispersion of traditional statistics is that in a sample of zeros, the measures of dispersion in traditional statistics are equal to zero. However, the empirical sample statistics of the Second Method in a sample of zeros would be equal to the inversion of N , the square of the inversion of N , and the square root of the square of the inversion of N .

The reason why, in a sample of zeros, the Mean Deviation of the Second Method is equal to the inversion of N is because the inversion of N is a theoretical probability. Regardless of the fact that the sample is a sample of zeros, the theoretical probability is a constant equal to the inversion of N . Even in a sample of zeros, theoretically, the theoretical probability of any random probability continues to be equal to the theoretical probability equal to the inversion of N .

Although the Bias Level is similar to the differential score, in reality they are not the same. The differential score in traditional statistics is only individual dispersion, but, in reality, the normal Bias Level fulfils a double function: the normal Bias Level measures the individual dispersion of the subject or option with respect to the norm (arithmetic mean equal to the inversion of N), but it also measures the Bias Level of the subject or option in relation to what would be the theoretical random behaviour ($1/N$).

Thus, if the Mean Deviation of the Second Method is equal to the inversion of N even in a sample of zeros, the Variance of the Second Method in a sample of zeros is equal to the square of the inversion of N ; therefore, the Standard Deviation is equal to the square root of the square of the inversion of N .

Because in Impossible Probability the value of the differentials in sample studies, in itself, is more important than the consideration of the signs, in Impossible Probability the Mean Deviation is more important than the Variance or the Standard Deviation, as it preserves the most reliable value of the differentials.

12. Theoretical Sample Statistics

If the empirical sample statistics are the Mean Deviation, the Variance and the Standard Deviation, which are in essence the average of the differential scores, then the theoretical sample statistics must be, in essence, the average of the maximum possible theoretical biases.

In the Second Method, the Maximum Possible Positive Bias is equal to one minus the inversion of N , and it only occurs when, out of the entire sample N , only one subject or option is equal to one; then, all other subjects or options are equal to zero, zero being the Minimum Possible Theoretical Probability, and the Maximum Possible Negative Bias is equal to zero minus the inversion of N .

Maximum Possible Positive Bias = $1 - 1/N$

Maximum Possible Negative Bias = $0 - 1/N$

In Impossible Probability it is said that the conditions of maximum dispersion only occur when, out of the entire sample N , only one subject or option is different from zero; then, the empirical probability of that single subject or option different from zero must be one. Thus, if only one subject or option is different from zero, all other subjects or options are equal to zero, that is, $N - 1$ subjects or options have probability zero.

If these conditions occur, then only one subject or option would have a Bias Level equal to one minus the inversion of N , while the rest of the subjects or options, $N - 1$, would have a Bias Level equal to zero minus the inversion of N , which in absolute terms, for those subjects or options equal to zero, would be the inversion of N .

Under conditions of maximum bias —a single subject or option equal to one, all remaining subjects or options equal to zero—, we can then calculate the Maximum Possible Theoretical Mean Deviation, equal to the average of the sum of the Maximum Possible Bias plus the product of the inversion of N by N – 1. And if we can calculate the Maximum Possible Theoretical Mean Deviation, then we can calculate the Maximum Possible Theoretical Variance and the Maximum Possible Theoretical Standard Deviation.

$$\text{Maximum Possible Theoretical Mean Deviation} = \{(1 - 1/N) + [(N - 1) \cdot 1/N]\} / N$$

$$\text{Maximum Possible Theoretical Variance} = \{(1 - 1/N)^2 + [(N - 1) \cdot (1/N)^2]\} / N$$

$$\text{Maximum Possible Theoretical Standard Deviation} = \sqrt{\{(1 - 1/N)^2 + [(N - 1) \cdot (1/N)^2]\} / N}$$

In the particular case of the Maximum Possible Theoretical Mean Deviation, alternative formulation models may also express exactly the same concept. Due to natural statistical goodness, it turns out that the Maximum Possible Theoretical Bias is equal to the product of the Maximum Possible Negative Bias by N – 1.

$$(1 - 1/N) = [(N - 1) \cdot 1/N]$$

The double of the Maximum Possible Positive Bias is equal to the sum of the Maximum Possible Positive Bias plus the product of the Maximum Possible Negative Bias multiplied by N – 1. And in the same way the double of the product of the Maximum Possible Negative Bias multiplied by N – 1 is equal to the sum of the Maximum Possible Positive Bias plus the product of the Maximum Possible Negative Bias multiplied by N – 1.

$$(1 - 1/N) + [(N - 1) \cdot 1/N] = 2 \cdot (1 - 1/N) = 2 \cdot [(N - 1) \cdot 1/N]$$

Thus, the Maximum Possible Theoretical Mean Deviation can also be expressed as the average of the double of the Maximum Possible Positive Bias.

$$\text{Maximum Possible Theoretical Mean Deviation} = [(1 - 1/N) \cdot 2] / N$$

13.Normal Bias Level and Object of Study

In traditional statistics, all statistical study, both in descriptive statistics and in inferential statistics, revolves around the dialectic between dispersion and central tendency. In Impossible Probability this dialectical tension is resolved in the identity of opposites, whose highest paradigm is the inversion of N itself.

It is very difficult to maintain the distinction between central tendency and dispersion when the inversion of N is, at the same time, the arithmetic mean and the theoretical probability of dispersion. In reality, in the Second Method there is no empirical arithmetic mean; the only thing we have is a theoretical estimate, the inversion of N, which serves multiple functions: mean, randomness and equality, and also theoretical dispersion probability in infinite universes, and probability of sampling-representativity error in infinite universes.

Due to the complexity of the multiple functions assumed by the inversion of N, in reality, in the Second Method of Impossible Probability, for data analysis there is nothing similar to what might be considered the empirical arithmetic mean.

Because of the absence of an empirical arithmetic mean —since the inversion of N is, in reality, a theoretical construct— although the Bias Level of the Second Method and the differential score in traditional statistics share similarities, they are in fact not the same, since they measure different qualities.

The differential score of traditional statistics measures the distance between a subject's raw score and the empirical arithmetic mean.

The Bias Level of a subject or option measures the distance between empirical probability and theoretical probability. It is not the same to measure the empirical distance between an empirical score and the empirical mean as to measure the distance between an empirical probability and a theoretical probability. Both distances may be interpreted as dispersion, we do not deny this, but dispersion relative to different reference points.

In traditional statistics, dispersion is between two empirical values: the raw score and the mean of the raw scores.

In the Second Method, dispersion is between an empirical value, the empirical probability, and a theoretical value, the inversion of N.

While traditional statistics measures dispersion between empirical values, in reality the Second Method measures dispersion between empirical values and a theoretical constant value, which is the inversion of N.

The importance of this theoretical value in the Second Method lies in the fact that it is not an arithmetic mean; the inversion of N is the theoretical probability that every subject or option should have at random in the absence of bias, under conditions of equal opportunity.

That is, if all students enjoyed the same conditions —good diet, good hygiene, regular rest, and adequate distribution of study time— all students, theoretically, should have the same academic outcome. Then, their empirical probabilities would be equal to the theoretical probability.

In reality, what the Bias Level measures is the distance between actual behaviour and theoretical random behaviour under equal opportunity and absence of bias. And it is for this reason that, even in a sample of zeros, the Mean Deviation is not equal to zero, because in fact it is the Mean Deviation between actual behaviour and random behaviour.

Because, while similar in form, the Bias Level is not the same as the differential score, this is the reason why the interpretation of the normal Bias Level will depend on the object of study.

In normal models, where the statistical norm is the inversion of N, the normal Bias Level can only tend towards three possible outcomes: positive bias, negative bias, or a tendency toward zero.

A subject or option is said to have positive bias when empirical probability exceeds the inversion of N. A subject or option is said to have negative bias when empirical probability is lower than the inversion of N. A subject or option is said to tend toward equality of

opportunity when, in absolute terms, the difference between empirical probability and the inversion of N tends to zero.

Thus, in normal studies (as opposed to omega models), there are three possible objects of study:

- studies of positive bias, when the ideal is to maximise the empirical probability of a single subject or option in the entire sample (conditions of maximum dispersion);
- studies of negative bias, when the ideal is to reduce empirical probability to zero (whether because $N - 1$ subjects or options tend toward zero, or because the ideal is the sample of zeros, or any other case in which we need to maximise negative bias);
- and equality studies, when the ideal is that every empirical probability of every subject or option should tend toward the inversion of N .

14. Omega Models

Normal models are those where the statistical norm is the arithmetic mean; in the case of Impossible Probability, the statistical norm is the inversion of N .

In contrast to the normal models based on the inversion of N , different non-normal statistical models could be constructed based on different statistical referents. Proof of this can be found in the so-called relative statistics, where, unlike normal statistics, the statistical referent is not the norm—the arithmetic mean—but any probability that, for a specific reason, may assume the role of statistical referent.

In this way, we can construct the statistic relative to the maximum, $p(x_i+)$, from the differentials between the empirical probabilities of a particular sample and the maximum empirical probability of that particular sample. The same can be applied to the minimum empirical probability $p(x_i-)$ of that sample, or to the intermediate empirical probability $p(x_i+/-)$, or to the empirical probability closest to the inversion of N , $p(x_i\approx)$, or we could construct a statistic relative to any empirical probability $p(x_n)$ of that particular sample.

Ultimately, the entire traditional construction of statistics begins from the consideration of the arithmetic mean as the central axis of the statistical study, which follows from the fact that traditional statistics is anchored in the axis dispersion vs. central tendency. But once we dissolve this tension into dialectical identity and discover that another form of statistics is indeed possible, we realise that the limits of traditional statistics could be surpassed at the very moment we understand that the barrier of the statistical norm can be overcome by accepting the possibility of different statistical referents, and not only the arithmetic mean.

With this, we are not denying the importance of traditional statistics, which we always refer to as the first method of reference; what we are trying to say is that it is possible to create alternative statistical models that may offer significant contributions to the development of statistics.

While in the textbooks of secondary schools and universities statistics and probability are presented as a completed science, with nothing more to be said, what we are demonstrating in *Impossible Probability* is that, even in terms of descriptive statistics, there are many advances that could be made simply by dismantling the myth that mathematics are what they are and there is nothing more to invent.

What we really want to say is that mathematics, far from being finished, are in fact yet to be created. In our opinion, the mathematics of today, 2 December 2025, will be very different from the possible mathematics we may have in 100 or 200 years, or why not, within the next 2000 years, thanks to the contribution of Artificial Intelligence. This is why we deeply believe that, at the moment in which Global Artificial Intelligence materialises—although probably with a more complex architecture and a different name—such intelligence will be capable of creating new mathematical models that surpass human ones; in a certain sense, such intelligence will produce non-human mathematics.

Once we understand that normal statistics, as those based on the statistical norm, although representing the first method of analysis, could be complemented with new statistical models, we can understand the importance of relative statistics.

In reality, what we refer to as omega models would, in fact, be a relative statistic, understood as those statistics that do not follow the statistical norm.

In the Second Method of Impossible Probability, omega models are those models in which, given a set N of subjects or options, there exists a subset of subjects or options greater than one and smaller than N, which we call omega, Ω , and the object of study is that all subjects or options within omega have probability different from zero without any hierarchical organisation; that is, only and exclusively the subjects or options of the subset omega within N will have empirical probability different from zero, and all subjects or options of omega will tend toward the same probability different from zero, the inversion of omega, which would be the ideal probability for every subject or option belonging to the omega subset.

Ideal omega probability = $1/\Omega$

Because in omega models the ideal probability is the inversion of omega, the inversion of omega becomes the statistical referent, so that the Bias Level does not follow the normal model but is studied as the difference between the empirical probabilities within omega and the ideal probability, which will be called the ideal Bias Level or omega Bias Level.

ideal Bias Level= omega Bias Level = $p(x_i) - 1/\Omega$

If the omega or ideal Bias Level is taken as the statistical referent, then the ideal or omega Bias Level of all subjects or options in omega would be equal to zero, and the ideal or omega Bias Level of all non-ideal subjects or options would be equal to the inversion of omega in absolute terms; thus, the Theoretical Mean Deviation Relative to Omega would be equal to the average of the product of the inversion of omega multiplied by the number of subjects or options not belonging to omega.

Theoretical Mean Deviation Relative to Omega = $(1/\Omega \cdot (N - \Omega)) : N$

The Theoretical Mean Deviation Relative to Omega is fulfilled only under the condition that all ideal subjects or options within omega are different from zero, and all non-ideal subjects or options outside omega are equal to zero.

In reality, for the Theoretical Mean Deviation Relative to Omega to be fulfilled, it is only necessary that all non-ideal subjects or options be equal to zero, which does not imply that all ideal subjects or options must necessarily be equal to the inversion of omega.

It may happen that all non-ideal subjects or options are equal to zero, and that only and exclusively the ideal subjects or options are non-zero. Even if all are non-zero, they may not all be equal to the inversion of omega, meaning divergences may exist among them, where there are ideal subjects or options different from zero but greater than the inversion of omega, or lower than the inversion of omega, even while all non-ideal subjects or options remain equal to zero.

For this reason, more important in critical omega sample studies will be the need to calculate the Ideal Mean Deviation, as the one that estimates the differences of all subjects or options, omega and non-omega, in relation to the inversion of N.

$$\text{Ideal Mean Deviation} = [(1/\Omega - 1/N) \cdot \Omega] + (1/N \cdot (N - \Omega)) / N$$

$$\text{Ideal Variance} = [(1/\Omega - 1/N)^2 \cdot \Omega] + ((1/N)^2 \cdot (N - \Omega)) / N$$

$$\text{Ideal Standard Deviation} = \sqrt{[(1/\Omega - 1/N)^2 \cdot \Omega] + ((1/N)^2 \cdot (N - \Omega)) / N}$$

15. Critical Probability and Inferential Statistics

Up to now, everything we have analysed falls within descriptive statistics—statistics that is solely dedicated to describing the sample using the most representative sample statistics.

In traditional statistics, the first method of analysis—regardless of whether the Second Method of analysis is adopted—revolves around descriptive statistics of central tendency and dispersion.

In the first data analysis we normally focus on the central tendency of raw scores or frequencies, analysing the arithmetic mean of the raw scores, the mode (the raw score with the highest frequency), and the median (the subject located around the 50th percentile of the statistical distribution). Among the dispersion statistics, the Mean Deviation of the raw scores, the Variance, and the Standard Deviation will be analysed.

In a second data analysis the Second Method may be applied, and in that case what must be identified is the empirical probability of each subject or option, the inversion of N, and the empirical and theoretical individual and sample statistics.

Descriptive statistics will only provide a description of the sample. In terms of Global Artificial Intelligence, descriptive statistics may have value in the application of inductive methodology for the creation of new categories in the categorical application. If in a data collection about a certain new object, the behaviour of the new object described in that collection cannot be attributed to any category in the application, then the application should automatically assume the data collection of the new object as a new category to be included in the application for the newly discovered object.

This process of scientific discovery through empirical induction, insofar as a scientific discovery has been made without having deduced a hypothesis a priori, will be a typical process in Specific Artificial Intelligence by Application for incorporating new categories into the application.

At the moment in which Artificial Research by Application can make the leap from Specific Artificial Intelligence by Application to the Unified Application, and to the synthesis of Unified Application and Global Artificial Intelligence, finally the Unified Application will proceed to incorporate new categories within the Global Artificial Intelligence through empirical induction.

Unlike global/specific deductive programmes, which will contribute new rational hypotheses derived from the contrast of empirical hypotheses.

While descriptive statistics can be applied to scientific induction, inferential statistics is the basis of the deductive method, which will give rise to Artificial Research by Deduction, first in Specific Artificial Intelligences by Deduction and finally evolving into the second stage of the Artificial Global Intelligence.

The main difference between description and inference is that description cannot provide new knowledge except through empirical induction (empiricism), whereas inferential statistics is the foundation of rational deduction (critical rationalism).

Although empirical induction has been heavily criticised in the 20th century, especially after the failure of the Vienna Circle to find any valid criterion of empirical meaning applicable to all scientific propositions, our view is that empirical induction still has great value in exploratory work.

Suppose we send a space mission to an unknown planet. Since we have no prior knowledge of that planet, it is impossible to make a priori deductions; therefore, no prior empirical hypothesis is possible. All the knowledge obtained from the initial exploration of that planet will be purely descriptive, meaning that all the knowledge obtained by the space mission about that planet will be knowledge acquired through empirical induction, the establishment of categories that define and describe that planet, and this will be the foundation of future applications designed specifically for that planet.

Once the first space mission has generated an application for that planet based on the categories and evidence found, empirical hypotheses can be formed that can be contrasted in future space missions.

Those future space missions will have the purpose of testing these hypotheses; therefore, they will be space missions whose source of knowledge will properly be inferential statistics.

But to reach that point it was necessary first to have clear descriptive knowledge and the establishment of the first applications that will support the application for that new planet.

To carry out the contrast of empirical hypotheses it will be necessary to establish, along with the hypotheses, rational criteria for the rational critique of the hypotheses.

The rational critique of the hypotheses will be made by producing collections of empirical data whose results will be contrasted with the rational criteria in order to demonstrate whether the hypotheses are correct.

These rational criteria will be called critical probabilities, and in Impossible Probability they are symbolised as " $p(xc)$ ". Since every probability is a quotient, a rational number, a ratio, the critical probability is essentially a critical rational number: the critical reason.

The way in which the critical ratio is established for hypothesis testing will be by multiplying a probability X , of error or reliability, by a theoretical statistic. That probability X of error or reliability is the percentage of error or reliability we are willing to accept with respect to the theoretical statistic.

The probability X of error or reliability is actually a percentage, but since we work with probabilities, not percentages, the percentage will always be presented divided by 100, adopting the form of a probability.

Thus, the probability X of error or reliability cannot be a reliability percentage greater than one nor an error percentage lower than zero.

X = probability of error or reliability

$X = 1 \rightarrow$ Maximum possible critical reliability

$X = 0 \rightarrow$ Minimum possible critical error

The probability X of error or reliability—1 maximum possible reliability, 0 maximum possible error—will be multiplied by the theoretical individual or sample statistic, depending on what is being critiqued.

If we are critiquing individual empirical statistics, then the probability X of error or reliability must be multiplied by the corresponding individual theoretical statistic.

If we are critiquing empirical sample statistics, then the probability X of error or reliability must be multiplied by the corresponding theoretical sample statistic.

The way in which the critique of reality is carried out and the way in which results are interpreted will depend on the object of study.

- If the object of study is equality of opportunity, then the absolute value of the individual normal Bias Level should be similar to the theoretical probability, within the margin of error we are willing to accept: the critical probability equal to probability X of error multiplied by the inversion of N .

At the sample level, in equality studies, the Mean Deviation should be equal to or less than the critical probability of error we are willing to accept, which would be the critical probability understood as probability X of error multiplied by the

Maximum Possible Theoretical Mean Deviation. The closer X is to zero, the less error we are willing to accept.

- In studies of individual positive bias, the normal Bias Level of the individual subject or option must be equal to or greater than a critical reliability probability: the critical probability will be equal to probability X of reliability multiplied by the Maximum Possible Theoretical Positive Bias.
- In studies of individual negative bias, the absolute value of the normal Bias Level must be equal to or greater than a probability X of reliability multiplied by the inversion of N .

In both studies of positive and negative bias, insofar as the object is to increase dispersion, at the sample level the critique will be carried out by analysing whether the Mean Deviation is equal to or greater than a critical value equal to the Maximum Possible Theoretical Mean Deviation multiplied by a probability X of reliability.

Rational critique is done at two levels—individual and sample—in order to ensure that the results are reliable.

As mentioned previously, the reasons why results may not be reliable are due to biases introduced by the sampling technique or due to the sample size itself, with sample size being a potentially even more important bias than the sampling technique.

By carrying out a double critique we can guarantee greater reliability of results, although obviously human technology is not free from error. In any case, our hope is that our technology will be able to resolve this problem in the near future.

In the following sections we will analyse the different techniques of critical analysis, validity and significance in individual rational critique, as well as sample rational critique by critiquing the Mean Deviation.

As previously mentioned, we prioritise the Mean Deviation because it preserves the absolute differentials without undergoing any transformation, whereas Variance

calculates the average of squared differentials, and Standard Deviation is the square root of the Variance.

If the research team considers that, in their judgement, the Variance or the Standard Deviation makes more sense for their investigation—normally because the first method always works with Variance or Standard Deviation—it is simply enough to replace in the calculation of the critical probability the Maximum Possible Theoretical Mean Deviation with the Maximum Possible Theoretical Variance or the Maximum Possible Theoretical Standard Deviation, and to contrast the critical values in relation to Variance and Standard Deviation

16. Critical Bias Level, Validity and Individual Significance

Once we have the possibility of establishing critical probabilities as statistical referents, it is possible to construct descriptive statistics relative to a given critical probability.

Insofar as a relative statistic results from comparing the probabilities of a sample in relation to any statistical referent different from the norm, we can establish statistics relative to the maximum empirical value of the sample, the minimum, the intermediate, etc., this being, in essence, the basis for the establishment of the Omega Models.

In the same way, relative statistics could be constructed for a given critical probability, simply by comparing the empirical probabilities in relation to a critical probability. This comparison would be called directly Critical Bias Level, as opposed to the normal Bias Level based on the inversion of N, or the omega Bias Level based on the inversion of omega. In the Critical Bias Level, one directly compares an empirical probability $p(x_i)$ and a critical probability $p(x_c)$.

Critical Bias Level = $[p(x_i) - p(x_c)]$

In the Critical Bias Level, the critical probability is directly the percentage X of error or reliability determined by scientific policy divided by one hundred; that is, the critical probability is the probability X of error or reliability determined by the scientific team.

In the Critical Bias Level, if the object of study is positive bias, the critical probability will be a reliability probability as close as possible to one; therefore, the margin of error would be equal to the difference of one minus the critical probability. In the Positive Critical Bias Level, one directly compares the empirical probability with the critical reliability probability; if the result is zero or positive, the tendency toward positive bias of the subject or option would be accepted.

Positive Critical Bias Level = $p(x_i) - p(x_c)$ = zero or positive $\rightarrow$ positive bias is accepted

$$p(x_c) = X : 100$$

X = reliability percentage

In studies of negative bias, the critical probability X would be a critical error probability, and if the difference between the critical error probability and the empirical probability is equal to zero or positive, the tendency of the subject or option toward negative bias would be accepted.

Negative Critical Bias Level = $p(x_c) - p(x_i)$ = zero or positive $\rightarrow$ negative bias is accepted

$$p(x_c) = X : 100$$

X = error percentage

The limitation of the Critical Bias Level in individual critique is that it can only critique positive bias or negative bias; it cannot be applied to individual critique of equality.

Furthermore, the establishment of the critical value —the critical probability— is made *ex nihilo*, that is, it is not obtained by calculating the percentage of error or reliability over an individual theoretical statistic, which makes the Critical Bias Level an individual critique system for positive or negative bias that lacks a theoretical referent.

Although, at an exploratory level, it is valid to compare empirical probability directly with a critical probability for positive or negative bias, and although it offers the advantage of enabling the development of critical relative descriptive statistics —in the sense that one can create a relative statistic based on a critical probability—, it would be advisable to combine this analysis with Validity and Significance of Positive or Negative Bias. In individual studies of equality, studies must be carried out directly using Validity of Equality and Significance of Equality.

While the individual Critical Bias Level is only valid for individual positive or negative bias, Validity and individual Significance are both valid for the critique of individual positive bias, individual negative bias, and individual tendency toward equality of opportunity. Their formulation is presented below:

Validity of Equality = $p(x_c) - / (p(x_i) - 1/N) / =$ zero or positive $\rightarrow$ equality of opportunity is accepted

$$p(x_c) = (1 - 1/N) \cdot (X : 100)$$

X = error percentage

Validity of Positive Bias = $(p(x_i) - 1/N) - p(x_c) =$ zero or positive $\rightarrow$ positive bias is accepted

$$p(x_c) = (1 - 1/N) \cdot (X : 100)$$

X = reliability percentage

Validity of Negative Bias = $(1/N - p(x_i)) - p(x_c) =$ zero or positive $\rightarrow$ negative bias is accepted

$$p(x_c) = 1/N \cdot (X : 100)$$

X = reliability percentage

Significance of Equality = $[(1 - 1/N) - / (p(x_i) - 1/N) /] - p(x_c) =$ zero or positive $\rightarrow$ equality is accepted

$$p(x_c) = (1 - 1/N) \cdot (X : 100)$$

X = reliability percentage

Significance of Positive Bias = $p(xc) - [(1 - 1/N) - (p(xi) - 1/N)] = \text{zero or positive} \rightarrow$
positive bias is accepted

$$p(xc) = (1 - 1/N) \cdot (X : 100)$$

X = error percentage

Significance of Negative Bias = $p(xc) - [1/N - (1/N - p(xi))] = \text{zero or positive} \rightarrow$ negative
bias is accepted

$$p(xc) = 1/N \cdot (X : 100)$$

X = error percentage

17. Critical Sample Level and Sample Significance

At the individual level, the Positive or Negative Critical Bias Level, Validity and Significance of Equality and Positive or Negative Bias, all of them are individual critique tests. The purpose of individual critique is to study whether the individual behaviour follows a tendency consistent with the empirical hypothesis according to its object of study: equality or bias, positive or negative.

While at the individual level we can obtain a rational individual verification of the hypothesis, in order to minimise any possible error—which may be due to the technique used or to the sample size—it is necessary to carry out a rational sample critique.

It is not enough to confirm that an individual behaviour is rational; it is necessary to demonstrate that the behaviour of the sample as a whole is rational. Once we achieve a double verification of rational behaviour at both the individual and collective levels, we can affirm the rational verification of the hypothesis, such that the empirical hypothesis becomes considered a rational hypothesis and becomes part of the rational truth.

In any case, a rational truth is not an absolutely pure truth; it is simply a linguistic proposition that provisionally reaches a validation of rational certainty.

The rational hypothesis is universal in every space in which the conditions under which it was demonstrated are met, and for the duration of its rational validity. The time of rational validity depends on the moment when any evidence, however minimal, is found that demonstrates its irrational character. The irrational character of a hypothesis is determined by the accepted margin of error.

To accept an error is always an error, but a necessary error as long as human intelligence maintains the physiological limits of its existence, an intelligence limited by its own cognitive, transcendental or genetic apriorisms. In any case, the acceptance of error is limited until the moment the error becomes visible; at that moment the rational hypothesis becomes demonstrably false, rejected, and ceases to be part of rational truth. These philosophical considerations inherent to Impossible Probability will be analysed in the second volume of this work.

For the validation of an empirical hypothesis at the sample level in the Second Method of Impossible Probability, two different concepts have been designed, both applicable to studies of bias and studies of equality: the Critical Sample Level of Bias or Equality, and the Sample Significance of Bias or Equality.

It should be noted that the Critical Sample Level of Bias and the Sample Significance of Bias are both valid for studies of negative bias as well as positive bias, insofar as in both cases the focus is on bias; therefore, the purpose is increasing dispersion, which implies increasing bias, and in both cases the Critical Sample Level of Bias or the Sample Significance of Bias may be applied.

In studies of equality, the Critical Sample Level of Equality or the Sample Significance of Equality will be applied.

In bias studies using the Critical Sample Level, what is examined is whether the sample dispersion, the Mean Deviation, is equal to or greater than a critical reliability probability estimated over a critical reliability probability multiplied by the Maximum Possible Theoretical Mean Deviation.

In equality studies using the Critical Sample Level, the opposite occurs: the object is to demonstrate that empirical sample dispersion, the Mean Deviation, is equal to or below a critical error probability estimated over an error probability multiplied by the Maximum Possible Theoretical Mean Deviation.

Studies of Sample Significance of Equality or Bias are simply the adaptation of individual significance studies to sample-level studies. Below, we present the formulations of the Critical Sample Level and Sample Significance.

Critical Sample Level of Equality = $p(xc) - DM = \text{zero or positive} \rightarrow \text{equality is accepted}$
 $DM = \text{Mean Deviation}$

$$p(xc) = \{ [(1 - 1/N) \cdot 2] : N \} \cdot (X : 100)$$

$X = \text{error percentage}$

Critical Sample Level of Bias = $DM - p(xc) = \text{zero or positive} \rightarrow \text{bias is accepted}$
 $DM = \text{Mean Deviation}$

$$p(xc) = \{ [(1 - 1/N) \cdot 2] : N \} \cdot (X : 100)$$

$X = \text{reliability percentage}$

Sample Significance of Bias = $p(xc) - \{ [(1 - 1/N) \cdot 2] : N \} - DM = \text{zero or positive} \rightarrow \text{bias is accepted}$

$DM = \text{Mean Deviation}$

$$p(xc) = \{ [(1 - 1/N) \cdot 2] : N \} \cdot (X : 100)$$

$X = \text{error percentage}$

Sample Significance of Equality = $\{ [(1 - 1/N) \cdot 2] : N \} - DM - p(xc) = \text{zero or positive} \rightarrow \text{equality of opportunity is accepted}$

$DM = \text{Mean Deviation}$

$$p(xc) = \{ [(1 - 1/N) \cdot 2] : N \} \cdot (X : 100)$$

$X = \text{reliability percentage}$

18. Individual and Sample Rational Critique in Omega Models

In omega models, rational critique will normally be directed toward the critique of the subjects or options within the omega subset. The study of negative bias in the non-omega subjects or options can be carried out, at the individual level, using techniques for studying negative bias.

At the sample level, whether the study focuses on the omega tendency of the omega subjects or options, or on the negative tendency of the non-omega subjects or options, in any case the sample critique will be directed toward critiquing sample dispersion, using as theoretical referent for critique the Ideal Mean Deviation.

As in normal models, the rational critique for individual omega models will first be presented, followed by the omega rational critique models for the sample.

Validity of Omega (individual) = $p(x_i) - p(x_c)$ = zero or positive $\rightarrow$ ideal is accepted
 $p(x_c) = 1/\Omega \cdot (X : 100)$
X = reliability percentage

Average Validity of Omega (only for omega subjects or options) =
 $\sum (p(x_i) - p(x_c)) : \Omega$ = positive $\rightarrow$ ideal is accepted
 $p(x_c) = 1/\Omega \cdot (X : 100)$
X = reliability percentage

Ideal Individual Critical Similarity = $p(x_c) - / \log (p(x_i) : 1/\Omega) /$ = zero or positive $\rightarrow$ accepted
 $p(x_c) = / \log (X : 100) /$
X = reliability percentage

Omega Sample Level = $DM - p(x_c)$ = zero or positive $\rightarrow$ omega is accepted at the sample level
DM = Mean Deviation
 $p(x_c) = \{ \{ [(1/\Omega - 1/N) \cdot \Omega] + [1/N \cdot (N - \Omega)] \} : N \} \cdot (X : 100)$
X = reliability percentage

$$\begin{aligned} &\text{Critical Ideal Sample Similarity in Mean Deviation} = \\ &p(xc) - / \text{Log} \{ DM : \{ [(1/\Omega - 1/N) \cdot \Omega] + [1/N \cdot (N - \Omega)] \} : N \} \} / = \text{zero or positive} \rightarrow \\ &\quad \text{accepted} \\ &\quad DM = \text{Mean Deviation} \\ &\quad p(xc) = / \text{Log} (X : 100) / \\ &\quad X = \text{reliability percentage} \end{aligned}$$

19. Individual Theoretical Critical Proportions

At the moment we can know what the maximum individual dispersion is thanks to the individual theoretical statistics, we can then know the proportion of the individual empirical Bias Level in relation to the Maximum Possible Theoretical Positive Bias in studies of equality and positive bias, and the proportion of the individual empirical Bias Level in relation to the Maximum Possible Theoretical Negative Bias.

From the calculation of the proportion between individual empirical dispersion and maximum individual theoretical dispersion, depending on the type of study, we can proceed to the rational critique of individual dispersion.

Individual Theoretical Critical Proportion of Equality

$$\begin{aligned} &\{ 1 - [/ (p(xi) - 1/N) / : (1 - 1/N)] \} - p(xc) = \text{zero or positive} \rightarrow \text{equality of opportunity is} \\ &\quad \text{accepted} \\ &\quad p(xc) = X : 100 \\ &\quad X = \text{reliability percentage} \end{aligned}$$

Individual Theoretical Critical Proportion of Positive Bias

$$\begin{aligned} &[(p(xi) - 1/N) : (1 - 1/N)] - p(xc) = \text{zero or positive} \rightarrow \text{positive bias is accepted} \\ &\quad p(xc) = (X : 100) \\ &\quad X = \text{reliability percentage} \end{aligned}$$

Individual Theoretical Critical Proportion of Negative Bias

$$[(1/N - p(x_i)) : 1/N] - p(x_c) = \text{zero or positive} \rightarrow \text{negative bias is accepted}$$

$$p(x_c) = X : 100$$

$$X = \text{reliability percentage}$$

20. Individual Empirical Critical Proportions

Just as we can calculate the proportion of individual empirical dispersion in relation to the maximum individual theoretical dispersion, depending on the object of study, if we have the statistic of maximum individual empirical dispersion —the Maximum Possible Empirical Bias, equal to the Total Bias divided by two= $[\sum / (p(x_i) - 1/N) / 2]$; understanding Total Bias= $\sum / (p(x_i) - 1/N) /$; as the absolute sum of the Bias Levels—, then we can calculate the empirical proportion between the individual Bias Level and the Maximum Possible Empirical Bias. And depending on the object of study, we proceed with the corresponding rational critique.

Individual Empirical Critical Proportion of Equality

$$\{1 - \{ / (p(x_i) - 1/N) / : [\sum / (p(x_i) - 1/N) / 2] \} \} - p(x_c) = \text{zero or positive} \rightarrow \text{equality of opportunity is accepted}$$

$$p(x_c) = X : 100$$

$$X = \text{reliability percentage}$$

Individual Empirical Critical Proportion of Positive Bias

$$\{ (p(x_i) - 1/N) : [\sum / (p(x_i) - 1/N) / 2] \} - p(x_c) = \text{zero or positive} \rightarrow \text{positive bias is accepted}$$

$$p(x_c) = X : 100$$

$$X = \text{reliability percentage}$$

Individual Empirical Critical Proportion of Negative Bias

$$\{ (1/N - p(x_i)) : [\sum / (p(x_i) - 1/N) / 2] \} - p(x_c) = \text{zero or positive} \rightarrow \text{negative bias is accepted}$$

$$p(x_c) = \{ 1/N : [\sum / (p(x_i) - 1/N) / 2] \} \cdot (X : 100)$$

$$X = \text{reliability percentage}$$

21. Theoretical Critical Sample Proportions

If at the individual level we can critique the proportion between individual empirical dispersion and individual theoretical dispersion, at the sample level we can do exactly the same by calculating the proportion between sample dispersion and the Maximum Possible Theoretical Mean Deviation, proceeding to the rational critique of the result depending on the object of study—bias or equality.

Theoretical Critical Sample Proportion of Equality in Mean Deviation

$$\{ 1 - \{ DM : \{ [(1 - 1/N) \cdot 2] : N \} \} \} - p(xc) = \text{zero or positive} \rightarrow \text{equality of opportunity is accepted}$$
$$p(xc) = X : 100$$

X = reliability percentage

Theoretical Critical Sample Proportion of Bias in Mean Deviation

$$[DM : \{ [(1 - 1/N) \cdot 2] : N \}] - p(xc) = \text{zero or positive} \rightarrow \text{bias is accepted}$$
$$p(xc) = X : 100$$

X = reliability percentage

22. Empirical Critical Sample Proportions

At the theoretical level, the maximum possible theoretical sample dispersion is the Maximum Possible Theoretical Mean Deviation, the Maximum Theoretical Variance, and the Maximum Possible Theoretical Standard Deviation. But when the Theoretical Intermediate Probability was analysed—which is the average of the sum of the Maximum (Theoretically) Possible Empirical Probability, one, and the Minimum (Theoretically) Possible Empirical Probability, zero—it was already mentioned that the main quality of the Theoretical Intermediate Probability is that it is the maximum possible empirical sample dispersion. That is, the Mean Deviation can absolutely never be greater than dividing by two the sum of one plus zero; that is, the Mean Deviation can only be equal to

or lower than one half. Zero point five is the maximum value of empirical sample dispersion.

Even though the Theoretical Intermediate Probability, 0.5, is the maximum empirical value that sample dispersion can reach, in reality this value can only be reached when the conditions of maximum dispersion occur in universes limited to two options. As N increases, even when the object of study is bias, the natural tendency of empirical statistics is to tend toward zero.

To resolve this dilemma in the critique of empirical critical sample proportions, what will be done is to compare, in the form of a quotient, the difference between the maximum empirical probability $p(x_i+)$ in the sample and the minimum empirical probability $p(x_i-)$ in the sample, in relation to zero point five, and critique the result of this differential with a critical probability appropriate to the object of study.

Empirical Critical Sample Proportion of Equality

$$\{ 1 - \{ [(p(x_i+) - p(x_i-)) : 2] : 0.5 \} \} - p(x_c) = \text{zero or positive} \rightarrow \text{equality is accepted}$$

$$p(x_c) = X : 100$$

X = reliability percentage

Empirical Critical Sample Proportion of Bias

$$\{ [(p(x_i+) - p(x_i-)) : 2] : 0.5 \} - p(x_c) = \text{zero or positive} \rightarrow \text{bias is accepted}$$

$$p(x_c) = X : 100$$

X = reliability percentage

23. Probability of Equality and Probability of Bias, Positive or Negative

Individual or theoretical proportions, empirical or sample-based, are proportions representing an individual empirical value, Bias Level, or a sample value, Mean Deviation, in relation to maximum empirical or theoretical statistics.

The theoretical individual proportion is the comparison in the form of a quotient of the Bias Level in relation to the possible maximum theoretical individual biases. The

empirical individual proportion is the comparison in the form of a quotient of the Bias Level in relation to the quotient of the sum of all absolute biases divided by two.

At the sample level, the theoretical sample proportion is the comparison between the Mean Deviation and the Maximum Possible Theoretical Mean Deviation. And at the empirical level, the comparison between the result of dividing by two the difference between the maximum and minimum empirical sample probabilities, divided by zero point five.

In the individual or sample proportions, empirical or theoretical, what is really done is a comparison in the form of a quotient between individual or sample empirical values in relation to maximum-tendency statistics, empirical or theoretical. Essentially, what individual or sample proportions measure is the proportion of a particular subject or option, or of the sample in general, in relation to a maximum-tendency statistic, according to the object of study.

However, once we have the Total Bias as the absolute sum of all Bias Levels, at the moment we divide the individual Bias Level by the Total Bias we are not only calculating a proportion; in reality, we are dividing a part of a whole by the whole to which that part belongs, and that ratio between a part and its corresponding whole is what we normally call empirical probability.

Thus, if an empirical probability is simply the division of a part by the whole to which it belongs, in reality when dividing the Bias Level by the absolute sum of all biases we are calculating the probability of that Bias Level, and that bias probability may be the object of study in both equality and bias, positive or negative.

Once we have calculated the empirical probability of equality or bias, positive or negative, we can start the process of critique. Here, the critical probability is calculated over zero point five (0.5). This is because the Maximum Empirical Bias Possible is the Total Bias divided by two, $[\sum (p(x_i) - 1/N)] : 2$. In other words, it is not possible to have a probability bias greater than half of the Total Bias, at least in normal statistics.

Critique of Probability of Equality

$$\{ 0.5 - \{ / (p(x_i) - 1/N) / : [\sum / (p(x_i) - 1/N) /] \} \} - p(x_c) = \text{zero or positive} \rightarrow \text{equality is accepted}$$
$$p(x_c) = 0.5 \cdot (X : 100)$$

X = reliability percentage

Critique of Probability of Positive Bias

$$\{ (p(x_i) - 1/N) : [\sum / (p(x_i) - 1/N) /] \} - p(x_c) = \text{zero or positive} \rightarrow \text{positive bias is accepted}$$
$$p(x_c) = 0.5 \cdot (X : 100)$$

X = reliability percentage

Rational Critique of Probability of Negative Bias

$$\{ (1/N - p(x_i)) : [\sum / (p(x_i) - 1/N) /] \} - p(x_c) = \text{zero or positive} \rightarrow \text{negative bias is accepted}$$
$$p(x_c) = \{ 1/N : [\sum / (p(x_i) - 1/N) /] \} \cdot (X : 100)$$

X = reliability percentage

24. Rational Critique of the Standard Score

In traditional statistics, the standard score is defined as the quotient of the differential score divided by the Standard Deviation. In Impossible Probability, we believe that the Mean Deviation is more important than the Standard Deviation because it preserves the differentials untouched as they appear in the sample, without any variation caused by the successive transformations of squaring and subsequent square-rooting applied in the Standard Deviation.

Although we firmly believe that the Mean Deviation is more faithful to the original data, we understand that traditionally the different hypothesis-testing models have been constructed using the Gaussian curve and the Standard Deviation. Therefore, it is very likely that, even if the Second Method is eventually used in scientific research —indeed, we are using it right now—, it is very probable that it will tend to adapt to traditional customs, such as the greater use of the Standard Deviation instead of the Mean Deviation.

In any case, in every statistical model in which we use the Mean Deviation and the Maximum Possible Theoretical Mean Deviation, the only thing necessary to adapt it to traditional statistics is simply to use the Standard Deviation and the Maximum Standard Deviation, as has always been mentioned.

Remaining faithful to our own principles, but always leaving open the possibility of adapting our scientific method to traditional customs, we will now present how the critical analysis of the Standard Score would look in the Second Method using Mean Deviation, DM, while indicating throughout the possibility of adapting it to the first method simply by replacing Mean Deviation with Standard Deviation and its corresponding maximum theoretical statistic.

The way we will proceed in critiquing the Standard Score in the Second Method is by calculating the quotient of the Bias Level divided by the Mean Deviation, MD, and the result will be critiqued in relation to a critical ratio depending on whether the object of study is equality or bias, positive or negative.

Rational Critique of the Standard Score for Equality

$p(xc) - [(p(xi) - 1/N) / : DM] = \text{zero or positive} \rightarrow \text{equality of opportunity is accepted}$

DM= Mean Deviation

$p(xc) = \{ (1 - 1/N) : \{ [(1 - 1/N) \cdot 2] : N \} \} \cdot (X : 100)$

X = error percentage

Rational Critique of the Standard Score for Positive Bias

$[(p(xi) - 1/N) : DM] - p(xc) = \text{zero or positive} \rightarrow \text{positive bias is accepted}$

DM= Mean Deviation

$p(xc) = \{ (1 - 1/N) : \{ [(1 - 1/N) \cdot 2] : N \} \} \cdot (X : 100)$

X = reliability percentage

Rational Critique of the Standard Score for Negative Bias

$[(1/N - p(xi)) : DM] - p(xc) = \text{zero or positive} \rightarrow \text{negative bias is accepted}$

$p(xc) = \{ 1/N : \{ [(1 - 1/N) \cdot 2] : N \} \} \cdot (X : 100)$

X = reliability percentage

25. Critical Proportions of the Normal Standard Score

The Standard Score can be understood as the Bias Level divided by the Maximum Possible Theoretical Mean Deviation, and it can be critiqued in relation to a critical probability. Now, the reason we can determine the critical probability is because we know the maximum tendency toward which the Standard Score can move.

If we know the maximum tendency toward which the Standard Score can move, then the quotient of the Standard Score divided by its maximum tendency will be the proportion of that maximum tendency represented by the Standard Score. This quotient of the Standard Score divided by its maximum tendency is called the proportion of the Standard Score over its maximum tendency, and this proportion can be analysed critically depending on whether the study concerns equality or bias, positive or negative.

Critical Proportion of Standard Score for Equality

$\{ 1 - \{ [(p(x_i) - 1/N) / : DM] : \{ (1 - 1/N) : \{ [(1 - 1/N) \cdot 2] : N \} \} \} \} - p(x_c) = \text{zero or positive}$
→ equality is accepted

DM= Mean Deviation

$p(x_c) = X : 100$

X = reliability percentage

Critical Proportion of Standard Score for Positive Bias

$\{ [(p(x_i) - 1/N) : DM] : \{ (1 - 1/N) : \{ [(1 - 1/N) \cdot 2] : N \} \} \} - p(x_c) = \text{zero or positive} \rightarrow$
positive bias is accepted

DM= Mean Deviation

$p(x_c) = X : 100$

X = reliability percentage

Critical Proportion of Standard Score for Negative Bias

$$\{ [(1/N - p(x_i)) : DM] : \{ 1/N : \{ [(1 - 1/N) \cdot 2] : N \} \} \} - p(x_c) = \text{zero or positive} \rightarrow \text{negative bias is accepted}$$

DM= Mean Deviation

$$p(x_c) = X : 100$$

X = reliability percentage

26. Rational Critique of the Standard Score under Omega Conditions

In omega models, the critique of the Standard Score is carried out differently, insofar as the omega maximum-tendency statistics are different. Essentially, the Standard Score is the same —Bias Level divided by Mean Deviation— but the rational critique changes because the critical probability is calculated on the basis of the omega maximum tendency; in this case, the reliability probability X multiplied by the result of the quotient obtained by dividing the Ideal Omega Bias (inversion of omega minus inversion of N) by the Ideal Mean Deviation.

In omega models, the Rational Critique of the Standard Score under omega conditions should be applied to the Standard Score of every omega subject or option, excluding non-omega subjects or options. Non-omega subjects or options should be studied using rational critique models for the Standard Score in negative bias.

Rational Critique of Omega Standard Score (DM)

$$[(p(x_i) - 1/N) : DM] - p(x_c) = \text{zero or positive} \rightarrow \text{positive bias is accepted}$$

DM= Mean Deviation

$$p(x_c) = \{ (1/\Omega - 1/N) : \{ [(1/\Omega - 1/N) \cdot \Omega] + [1/N \cdot (N - \Omega)] \} : N \} \} \cdot (X : 100)$$

X = reliability percentage

27. Critical Ideal Similarity of the Standard Score

Given that, in omega models, the object of study is that the omega subjects or options tend toward the inversion of omega, the aim is the similarity of the behaviour of the omega subjects or options to the inversion of omega.

In Impossible Probability, the Similarity Level is the base-ten logarithm of the quotient between two values; in the case of omega models, it is the base-ten logarithm of the quotient of the empirical probability of an omega subject or option divided by the inversion of omega.

And in the specific case of the omega Standard Score, what we will critique is the level of similarity between the Standard Score of an omega subject or option and what its ideal would be, with the Ideal Standard Score being equal to the quotient obtained by dividing the differential of inversion of omega minus inversion of N by the Ideal Mean Deviation.

Critical Ideal Similarity of the Standard Score using Ideal Mean Deviation

$$p(xc) - \{ / \text{Log} \{ [(p(xi) - 1/N) : DM] : \{ (1/\Omega - 1/N) : \{ [(1/\Omega - 1/N) \cdot \Omega] + [1/N \cdot (N - \Omega)] \} : N \} \} \} / \} = \text{zero or positive} \rightarrow \text{positive bias is accepted}$$

$$\begin{aligned} DM &= \text{Mean Deviation} \\ p(xc) &= / \text{Log} (X : 100) / \\ X &= \text{reliability percentage} \end{aligned}$$

28. Critique of the magnitude N and the effect of N

Up to now what has been analysed are theoretical and empirical descriptive statistics over a particular sample, and models of inferential rational critique over a specific data collection, a single sample; basically we have only done descriptive and inferential intra-sample statistics, that is, starting from data taken from a single sample.

The reality is that in scientific research the ideal is inter-measurement studies, and if possible inter-sample studies.

By intra-measurement study we define the study of the data collection obtained from a single and exclusive unique sample. When we only have one single sample, the defence of the thesis will be made on that single and exclusive sample that we have, given that we do not have any other.

Suppose that we discover a meteorite on Earth, or an asteroid in space, with some type of chemical component unknown until now. In the first place, through methods of scientific induction we will be able to establish the corresponding categories for that chemical component, and from those first categories we will be able to advance in the deduction of hypotheses such as: possible origin and formation of that chemical component, possible health effects, or potential medical or industrial uses.

But insofar as that single meteorite or asteroid is the only sample we have, the defence of any thesis about that single meteorite or asteroid will necessarily be an intra-measurement and intra-sample defence: we only have one sample, that meteorite or asteroid. Or in any case an inter-measurement study but still intra-sample, if we take different measurements of the same sample over a period of time in order to know the evolution of the sample over a determined period. At the moment in which we only have a single sample, the objective will be to take as many measurements as possible of that same sample, for as long a time as we can preserve the sample for the study.

At the moment in which we can have more samples, for example, if we can know the origin of that meteorite or asteroid and send space missions to the source, then we will be able to obtain more evidence, more samples, more proofs, and at that very moment in which we have more samples we will be able to do inter-measurement inter-sample studies.

While in an intra-measurement intra-sample study there are situations where we cannot select the magnitude of the sample according to our desire—for example, if we only have one meteorite or asteroid, the sample does not depend on us, but is what we have found—, the only mission is to preserve the essence of the only sample we have for as long as possible. On the other hand, at the moment in which we have a greater quantity of samples, then we can indeed create more important samples in which to include the originals, and increase the measurements and the samples.

At the moment in which we can do inter-measurement inter-sample studies is when we can do sample selection, and it is at that moment when we must begin to critique the magnitude of the sample, to know whether it is sufficiently representative. And at the moment in which we make inter-sample comparisons, we must understand what the effect of N is in the comparison of results belonging to different samples of different N.

Here what we find are two related but different phenomena. In the first place, the critique of the sample magnitude; in the second place, the calculation of the effect N when statistics from samples of different N are compared.

In relation to the critique of the sample magnitude, it must be noted that we must not only critique the magnitude of the sample N, particularly important in studies of subjects, but also study the magnitude of the sample of frequencies unless the object of study is the sample of zeros or the moderation of behaviour. And, in addition, critique the relationship between the sample N and the sample of raw scores or frequencies.

In the particular case of the frequency sample, what is the basis of option studies, in option studies an ideal may be the zero sample, for example if we have a universe of symptoms in a disease, this case the ideal is to reduce the symptoms to zero, therefore the ideal is zero frequency per symptom, therefore the ideal is a zero frequency sample. In the critique of the raw-score or frequency sample the object of study must be taken into account. When we focus on negative bias, the frequency sample may descend considerably, therefore as the frequencies descend it would be understood as evidence of the descent of negative biases. Otherwise, in cases of equality or bias, positive or negative, the ideal is to have a sample frequency as large as possible.

The way in which one can critique the magnitude N or the magnitude of the raw-score or frequency sample is critiquing the inversion of N and critiquing the inversion of the raw-score or frequency sample. And the way to critique the relationship of N and the raw-score or frequency sample is critiquing the difference relationship of the inversion of N minus the inversion of the raw-score or frequency sample.

Sample Sufficiency $N = p(xc) - 1/N = \text{zero or positive}$ the sample is accepted

$$p(xc) = X : 100$$

$X = \text{error percentage}$

Sample Sufficiency of $\sum x_i$, in studies of equality or positive bias

$$p(xc) - 1/\sum x_i = \text{zero or positive it is accepted}$$
$$p(xc) = X : 100$$

Critical Requirement of Normality, especially in bias studies

$$(1/N - 1/\sum x_i) - p(xc) = \text{zero or positive normality is accepted}$$
$$p(xc) = 1/N \cdot (X : 100)$$

$X = \text{reliability percentage}$

Requirement of Normality, in equality of opportunities or moderation of behaviour

$$1/N - 1/\sum x_i = \text{zero or positive normality is accepted}$$

Flexible Sample Sufficiency, especially for discrete categories

$$p(xc) - N/\sum x_i = \text{zero or positive it is accepted}$$
$$p(xc) = X : 100$$

$X = \text{error percentage}$

The critique of the sample magnitude is only possible when we have the option to choose samples and when we have as many samples as necessary in order to guarantee a study as isomorphic as possible. As mentioned previously, although in market studies or in social sciences this is always possible, in other fields of science we do not always have the opportunity to have all the samples necessary or desirable; we simply must adapt to the samples we have due to the rarity of the phenomenon, or to the technological difficulty in accessing the phenomenon, or to the novelty of the phenomenon.

It is at the limits of science where we can find true challenges, and it is at the limits of science, science at the limit, where the real debate arises about what we understand by

science and what science itself is. From the comfortable armchairs of academic universities it is very easy to define science sitting in an armchair calmly having a coffee while you prepare the PowerPoint for your students. From a space mission to an unknown planet the vision of science changes completely: science is surviving.

When science is reduced to surviving, any smallest detail matters; the smallest error is fatal, and the difference between life and death is only seconds or millimetres: the margin of error is practically zero. Such is science in places like NASA or DARPA.

At the moment in which we know the source of that meteorite or asteroid, or at the moment in which we can send more space missions to an unknown planet, we can take more samples, and at that moment we can make comparisons between different samples of different N ; therefore, at that moment we must begin to include the effect of N in the comparisons of results between different samples.

The importance of the effect of N lies in the fact that, independently of the object of study, under normal conditions what is usual is that, as N increases, the tendency to zero of empirical probabilities increases, in the same measure that dispersion tends to zero.

Thus, in the analysis of bias in inter-measurement inter-sample studies, absence of bias may be attributed to a phenomenon in its evolution in different samples in different measurements simply as a consequence of not having included the effect of N in the different measurements. If in a first sample the value N is greater than the value N of the second measurement, logically the empirical probabilities of the second measurement will tend to be greater than the empirical probabilities of the first measurement, simply as a consequence of the effect of N in the different samples. This phenomenon could lead to the attribution of bias when in reality it is a product of the effect of N .

In the same way, the error could be committed of attributing to a behaviour equality of opportunities merely by the fact of having increased the magnitude N in successive later measurements, which would lead to a reduction of dispersion simply as a consequence of the reduction of the sample size, not due to the behaviour itself of the phenomenon in the different measurements in the different samples.

At the moment in which we enjoy the privilege of being able to count on different samples and different measurements, independently of whether there are divergences in the sample size, inter-measurement inter-sample comparisons could be made simply taking into consideration the effect of N on samples of different magnitude N.

To carry out comparisons between samples of different N what we proposed in *Introducción a la Probabilidad Imposible* is the designation of a sample with the measurement adjective “m”, normally the sample “m” is the sample N that will be used as a reference to compare a set of samples “ $m_i = m_1, m_2, m_3... m_i$ ”. Thus, to the magnitude N of the sample “m” we will symbolise it “ N_m ”, and any sample of the set “ $m_i = m_1, m_2, m_3... m_i$ ” will be symbolised with the symbol “ N_{m_i} ”. Under this legend then the following propositions can be made in relation to the sample “m”:

$N_m \neq N_{m_i}$

N_m = total number of subjects or options in measurement “m”

$1/N_m$ = theoretical probability of measurement “m”

$p(x/m)$ = empirical probability in measurement “m”

Then for any sample of the set “ m_i ” the legend will be the following simply substituting by “ m_i ”:

$N_{m_i} \neq N_m$

N_{m_i} = total number of subjects or options in a measurement “ m_i ”

$1/N_{m_i}$ = theoretical probability of a measurement “ m_i ”

$p(x/m_i)$ = empirical probability in a measurement “ m_i ”

Once we have the following legend, the way to compare the empirical probability or statistic of sample “m” with the empirical probability or statistic of a sample belonging to “ m_i ” of different sample size, is applying the effect of N on any empirical probability or statistic of any sample of the set “ m_i ”. Insofar as the effect of N is inverse to the magnitude N—the greater the N, the greater the tendency to zero; the smaller the N, the lesser the tendency to zero—the effect of N will be equal to the quotient of dividing the magnitude “ N_{m_i} ” by the magnitude “ N_m ”, and the effect of N will be multiplied by any empirical probability or statistic of any sample of the set “ m_i ” to compare with the sample “m”.

Effect of $N = N_{mi} : N_m$

Effect of N on " $p(x_i/m_i)$ " = $p(x_i/m_i) \cdot (N_{mi} : N_m)$

Effect of N on " $1/N_{mi}$ " = $1/N_{mi} \cdot (N_{mi} : N_m)$

In the particular case of the Effect of N on theoretical probability, $1/N$, what can be easily verified is that, being N_{mi} different from N_m , at the moment in which the Effect of N is applied to the inversion of N_{mi} , then the difference of the inversion of N_{mi} minus the inversion of N_m is equal to zero, which shows that its magnitude is the same. Therefore, applying the Effect of N to any empirical probability or statistic of any sample of the set " m_i ", it is possible to compare these empirical probabilities or statistics of the set " m_i " in relation to the sample " m ".

$$\underline{[1/N_{mi} \cdot (N_{mi} : N_m)] - 1/N_m = 0}$$

The reason for placing the analysis of the effect of N at this moment of the Analysis of Impossible Probability is because in the following sections we will make the exposition of the rational models of critique of necessary and possible variations, which are the basis for the elaboration of projections and forecasts, which are in essence the basis of all model and project theory in Global Artificial Intelligence.

In this Analysis of Impossible Probability we are prioritising a summarised analysis of the principal contributions of *Introducción a la Probabilidad Imposible*. We understand that the magnitude of the work, number of pages in Spanish and number of volumes in English, as well as the language used, may be a limit for the diffusion of this work; in reality we never thought that we would come to publish it for mass diffusion in life. The true reason for this Analysis of Impossible Probability is simply to give a more simple and reduced version of the principal contributions, and not only for experts or judges, but for any interested reader or students.

For this reason, in the study of variations, projections and forecasts, we will simply limit ourselves to saying that when studies require inter-measurement inter-sample comparisons, one must be aware that the factor N can influence and can be overcome using the Effect of N , simply identifying which is the sample whose magnitude N will be used as the reference, assigning it the function of sample " m "; then, its value N becomes

the magnitude “Nm”. Thus, all the other samples to compare must be integrated in the set “mi” and their respective magnitudes N will come to be conceptualised as “Nmi”.

From that moment, the Effect of N will be calculated as the quotient Nmi divided by Nm and the Effect of N will be multiplied by any probability or statistic of any sample of the set “mi” in order to be compared in relation to the sample “m”, or, once all the samples “mi” have been standardised and transformed by the Effect of N, to make inter-sample comparisons between different samples “mi” among themselves once the Effect of N has been applied to each sample “mi”.

29.Critical necessary variation and possible variation

Necessary variation constitutes the fundamental distance that exists between the empirical probability of a subject or option and any other individual statistic or empirical probability (from the same or a different sample) with which it is compared. In this sense, the Bias Level, both normal and relative, expresses not only the bias between two factors, but also the necessary variability existing between them, provided that said Bias Level is interpreted in its absolute value or the factors are ordered in such a way that the greater is subtracted from the lesser or vice versa. In this way, the necessary variation is that quantity which, added to the lesser of the factors, reproduces the greater, or that which, subtracted from the greater, yields the lesser.

The Bias Level thus fulfils two essential functions. In the first place, it is a strict measure of bias between two factors. In the second place, when it is interpreted in absolute value, it becomes a direct measure of the necessary variation between the same factors. For this reason, when it is considered as necessary variability, the Bias Level relative to the maximum or the minimum with respect to the other, or vice versa, operates according to the following structure:

$$\begin{aligned} & / [p(x+) - p(x-)] / = \text{necessary variation between the maximum and the minimum} \\ & / [p(x+) - p(x-)] / + p(x_i-) = p(x_i+) \end{aligned}$$

This implies that the absolute value of the Bias Level relative between the maximum and the minimum represents exactly the distance separating both factors, so that, when added to the lesser, it produces the greater. Conversely:

$$\begin{aligned} & / [p(x+) - p(x-)] / = \text{necessary variation between the maximum and the minimum} \\ & p(x_i+) - / [p(x+) - p(x-)] / = p(x_i-) \end{aligned}$$

This means that the same necessary variation, subtracted from the greater, yields the lesser. Thus the symmetric structure of necessary variability is defined: the necessary decrease of the greater to equal the lesser and the necessary increase of the lesser to equal the greater constitute two equivalent expressions of the same probabilistic distance.

From this conceptual basis, necessary variation can be understood as the real distance that exists between two factors. This distance simultaneously measures how much the probability of the greater factor should decrease in order to be equal to the lesser and how much the probability of the lesser factor should increase in order to be equal to the greater. The necessary variability of the greater to match the lesser is, therefore, a quantifiable decrease; that of the lesser to match the greater, a quantifiable increase. Both correspond to the same magnitude.

From this, different necessary variations can be defined formally:

– In equality of opportunities, the normal Bias Level would also fulfil the function of being the necessary variation between an empirical probability and the inversion of N:

$$/ (p(x) - 1/N) /$$

– In studies of positive bias, the necessary variation between the Maximum Possible Empirical Probability and an empirical probability:

$$(1 - p(x_i))$$

– In studies of negative bias, the necessary variation between an empirical probability and the Minimum Possible Empirical Probability:

$$(p(x_i) - 0) = p(x_i)$$

As a consequence, every absolute value of a normal or relative Bias Level, or the difference between an empirical probability and an individual theoretical statistic, always represents the necessary variation that would allow the greater to decrease enough to equal the lesser or the lesser to increase as needed to equal the greater. Every equality between factors requires a necessary variation equivalent to the absolute value of their differential.

However, in empirical reality, the fact that something is necessary does not imply that it is fully possible. Statistics operates on samples, and every sample is a fragment of reality, a perceptual illusion. Consequently, in statistics it is not possible to work on the universe itself, but only on the margins of possibility that the sample (the mind) allows. Within these margins, the necessary may be possible only in part, either by chance or in a reliable way, depending on the empirical or theoretical conditions available. Absolute impossibility exists only in the infinite; in a sample, what exists are limited possibilities.

For this reason, Possible Variation will study what part of the necessary variation can actually be manifested in the sample, whether by chance or through a reliable regularity, on the critical or empirical plane, individual or sample-based. It is on this basis that the conceptual transition between necessary variation and possible variation is structured.

In Introduction to Impossible Probability all the formulations related to the contrast of possible variations are presented in full, both those that can occur by chance and those that can be established in a reliable way. In the present edition of *Analysis of Impossible Probability* we will focus exclusively on the study of reliable possible variations, since

these are the ones that allow a solid foundation of subsequent analytical processes and those that are essential for the elaboration of consistent rational models.

For those readers who wish to deepen their understanding of the complete analysis, including possible variations by chance and the full development of the corresponding contrasts, it is recommended to read section 16 of Introduction to Impossible Probability, where the complete set of formulations and applied methods is presented in detail.

The contrasts of possible variations that will be presented in *Analysis of Impossible Probability* will concern inter-sample possible variations with equal N. In the case that inter-sample studies were carried out between samples of different N it would be necessary to identify which sample would assume the function of sample “m”, and the set of samples that would be included in the set “mi”, so that the Effect of N would be applied to the empirical probabilities and statistics of the samples of the set “mi”.

Those readers who wish to deepen these contents in Introduction to Impossible Probability will see not only the contrasts of possible variations by chance developed; in addition, the application of the Effect of N is included in the studies of variations by chance or reliable between samples of different N.

Because the variation studies will be studies in which the variation of behaviour between different samples or measurements is studied, the legend that will be used to designate the statistics of each sample or measurement will be the following:

$p(x_i/1^a)$ = empirical probability first measurement

$p(x_i/2^a)$ = empirical probability second measurement

In the case that the first measurement and the second measurement have different N, then decide which sample plays the role of “m”, and which sample plays the role of “mi”, and multiply the Effect of N by the empirical probability “mi”, as well as, if necessary, multiply the Effect of N by any statistic of the sample “mi” that needs it.

It should be noted that, as in other critical contrast models, not only is the critique of individual variations necessary, but the critique of sample variations is also necessary, so that, if we have more than one measurement and the object of study is equality of opportunities, in successive measurements the Mean Deviation should tend to zero

sufficiently for this to be rational; otherwise, it would be rejected that the behaviour evolves favourably towards equality of opportunities in the sample as a whole, even if some subjects or options individually did show an evolutionary behaviour rational towards equality of opportunities.

The fact that a group of subjects or options conforms to the object of study does not mean that the whole sample conforms to the object of study; it simply means that this specific group of subjects or options evolves favourably, while this may not occur in the sample as a whole.

In bias studies exactly the same thing happens, so it will be necessary to study how the variations in the increase of the sample bias, Mean Deviation, are sufficiently rational in accordance with the criteria of scientific research, the critical probability.

In the studies of sample variation between two different samples, it should be indicated that the Mean Deviation of a first measurement and the Mean Deviation of a second measurement are compared, which will have the following legend:

DM/1^a = Mean Deviation of the first measurement
DM/2^a = Mean Deviation of the second measurement

In the case that the measurements have different magnitude N, apply the Effect of N to the Mean Deviation whose sample is designated “mi”. Below are presented the contrast formulations for critical individual and sample possible variations.

Contraste de Variación Posible Fiable Crítico Individual, en estudio de igualdad de oportunidades

$$\left[\frac{1}{p(x_i/1^a) - 1/N} - \frac{1}{p(x_i/2^a) - 1/N} \right] - \left[\frac{1}{p(x_i/1^a) - 1/N} \cdot p(x_c) \right] = \text{zero or positive}$$

equality of opportunities is accepted

$$p(x_c) = X : 100$$

X = reliability percentage

Contraste de Variación Posible Fiable Crítico Individual, en estudio de sesgo positivo

$$(p(x_i/2^a) - p(x_i/1^a)) - [(1 - p(x_i/1^a)) \cdot p(x_c)] = \text{zero or positive reliable increase of positive bias is accepted}$$

$$p(x_c) = X : 100$$

X = reliability percentage

Contraste de Variación Posible Fiable Crítico Individual, en estudio de sesgo negativo

$$(p(x_i/1^a) - p(x_i/2^a)) - [p(x_i/1^a) - 0] \cdot p(x_c) = \text{zero or positive reliable increase of negative bias is accepted}$$

$$p(x_c) = X : 100$$

X = reliability percentage

Contrastes de Variación Posible Fiable Crítica Muestral en estudio de sesgo

$$\{ [DM/2^a - DM/1^a] - \{ \{ [(1 - 1/N) \cdot 2] : N \} - DM/1^a \} \cdot p(x_c) \} = \text{zero or positive it is accepted in bias study}$$

$$p(x_c) = X : 100$$

X = reliability percentage

Contrastes de Variación Posible Fiable Crítica Muestral normal en estudio de igualdad de oportunidades

$$(DM/1^a - DM/2^a) - [DM/1^{a-m} \cdot p(x_c)] = \text{zero or positive reliable tendency to equality of opportunities is accepted}$$

$$p(x_c) = X : 100$$

X = reliability percentage

30.Types of empirical error and empirical reliability

Up to now, whenever any contrast has been made it has been done using rational critical criteria established by scientific policy. For the moment we have never used empirical criteria for hypothesis testing. In a certain sense, the tradition of using only critical reasons for contrast is due to the very rationalist tradition, which always prioritises rational values over empirical ones in the contrast of ideas, assuming that the rational critical values established by the researcher will be guided by the rational interest in love of truth and science; that is, the true interest of science itself is to know what is

happening, and for this reason we assume that the ethics of the researcher is always to use the most demanding and rigorous critical reason possible in order to demonstrate the truth of their hypothesis.

In the field of Global Artificial Intelligence, at the moment when scientific research by Artificial Intelligence becomes possible, it will be Artificial Intelligence itself, without the need for human mediation, that decides the rational criteria. This, in a certain sense, from our point of view, will increase the rigour of science.

Although rational criteria can be established, whether of human origin or established by Artificial Intelligence, it must be indicated that the tradition of rational critique of hypotheses only by critical reason could be made compatible with empirical contrast models that do not use a critical reason.

In essence, a critical probability is a critical ratio; it is the way in which we materialise rational philosophy in mathematics: critical reason is materialised in scientific research in the form of a mathematical critical ratio, the critical probability.

In the tradition of critical rationalism, in which Impossible Probability is immersed, reality is traditionally criticised by critical reason, which is operated mathematically in the form of critical probability.

Now then, like every tradition, it is susceptible to changes and improvements. What we introduce in this section, on empirical error and reliability, and in the following section, on models of empirical critical contrast of individual and sample variation —and which is a faithful reflection of what can also be found in section 16 of Introduction to Impossible Probability— is that empirical error or reliability models can be used for the empirical critique of empirical hypotheses.

In this way, two different and compatible types of critique can be distinguished: rational critique in the form of critical probability, and empirical critique in the form of empirical criteria for hypothesis validation.

If in rational critique the critical reason is materialised in the form of the product of a percentage X of error or reliability multiplied by theoretical statistics, normally of maximum tendency, in the same way the empirical criteria of contrast will be criteria of

empirical error or reliability multiplied by theoretical statistics, normally of maximum tendency.

The only difference between a critical probability and an empirical contrast criterion is that the theoretical statistic in critical probability is multiplied by a percentage X of error or reliability, while in empirical contrast criteria the theoretical statistics are multiplied by empirical criteria of error or reliability.

At this point, the only thing we need in order to establish empirical criteria for hypothesis contrast is to locate the empirical error or reliability criteria.

The empirical error or reliability criteria in the Second Method depend on the object of study. In equality of opportunities, dispersion is understood as a source of error, insofar as the greater the dispersion, the lower the equality; therefore, dispersion is error. However, in bias models the opposite occurs: if the object of study is bias, then bias is reliability; therefore, the greater the dispersion, the greater the reliability.

If in equality dispersion is error, and in bias dispersion is reliability, then the way to calculate error in equality or reliability in bias is by dividing the Mean Deviation by the Maximum Possible Theoretical Mean Deviation. In other words, the proportion of Mean Deviation in relation to the Maximum Possible Theoretical Mean Deviation is what will determine the degree of error in equality or reliability in bias.

In this way, the first modality of empirical error will be called alpha, α , and is applied to studies of equality of opportunities; it will be equal to the division of the Mean Deviation by the Maximum Possible Theoretical Mean Deviation. In this way, in equality of opportunities, the first modality of reliability will be equal to the difference of unity minus alpha, $1 - \alpha$, and will be symbolised $\neg\alpha$.

Inversely, in bias studies, the second modality of empirical reliability will be equal to the quotient of dividing the Mean Deviation by the Maximum Possible Theoretical Mean Deviation, and will be called beta, β . Then, in bias studies, the second modality of empirical error will be equal to unity minus beta, $1 - \beta$, and will be represented $\neg\beta$.

It is mentioned again that, to adapt these formulations to the first method, traditional statistics, it is enough to substitute Mean Deviation for Standard Deviation, and to

substitute the Maximum Possible Theoretical Mean Deviation for the Maximum Possible Theoretical Standard Deviation. In the following exposition of formulations, given that in Impossible Probability we will always prioritise the Mean Deviation, we present the formulas with Mean Deviation. In any case, in *Introduction to Impossible Probability* they were also indicated using Standard Deviation. It is again recalled that the initials DM correspond to Mean Deviation.

$\acute{\alpha}$ = First Modality of Empirical Error

First Modality of Empirical Error, $\acute{\alpha}$, in DM, in equality of opportunities

$$\acute{\alpha} = DM : \{ [(1 - 1/N) \cdot 2] : N \}$$

$\neg \acute{\alpha} = 1 - \acute{\alpha}$ = First Modality of Empirical Reliability

First Modality of Empirical Reliability, $\neg \acute{\alpha}$, in DM, in study of equality

$$\neg \acute{\alpha} = 1 - \{ DM : \{ [(1 - 1/N) \cdot 2] : N \} \}$$

β = Second Modality of Empirical Reliability

Second Modality of Empirical Reliability, β , in DM, in bias study

$$\beta = DM : \{ [(1 - 1/N) \cdot 2] : N \}$$

$\neg \beta = 1 - \beta$ = Second Modality of Empirical Error

Second Modality of Empirical Error, $\neg \beta$, in DM, in bias study

$$\neg \beta = 1 - \{ DM : \{ [(1 - 1/N) \cdot 2] : N \} \}$$

31.The necessary variation and empirical possible variation

At the moment we have empirical criteria for contrast, we simply must apply the empirical contrast criteria according to the object of study (equality or bias, positive or negative) at both the individual and sample level.

The modalities of empirical error —first modality of empirical error in equality of opportunities, second modality of empirical error in bias studies—, in the specific case

of studies of variations, are applied to the analysis of studies by chance, that is, determining when a variation is due to chance.

Although in an in-depth study these types of concepts are important, given that in Analysis of Impossible Probability we focus on the most fundamental aspects to highlight (the objective is a clear and simple work on the most representative concepts of Impossible Probability), we will go directly to the analysis of empirical reliability contrasts applied to possible variations.

In any case, it should be emphasised that the fact that a phenomenon is not attributed to an experimental variable does not imply that it is due to chance, which is the reason why in Introduction to Impossible Probability both analyses are included, because there may be phenomena which, although not the product of our intervention, also cannot be attributed to chance, and there may exist other sources that intervene in the processes. When a phenomenon is not random, but neither is it the product of our intervention in reality, it requires further investigation to analyse why something that is not the result of our experimental manipulation nevertheless does not follow random patterns, which means there is an external intervention different from ours by factors we have not identified.

If a variation were to occur that is not attributable to our manipulation and not attributable to chance, the phenomenon requires an in-depth investigation to determine the source of the interferences.

For this reason, the analysis of when a phenomenon is by chance is just as important as the analysis of when a phenomenon is the product of experimental manipulation. Both analyses are essential; however, we repeat, the purpose of the present work is to offer the most essential concepts so that any reader, in a simple reading, may have at their disposal the most characteristic concepts of our theory.

Having said this, we proceed to the exposition of the empirical contrast models of equality or bias, positive or negative, at the individual and sample level, applied to possible variations, for which the first modality of empirical reliability in equality studies and the second modality of empirical reliability in bias studies are required.

Contraste de Variación Posible Fiable Empírica Individual, mediante Primera Modalidad de Fiabilidad Empírica, en igualdad de oportunidades

$$\{ / (p(xi/1^a) - 1/N) / - / [p(xi/2^a) - 1/N] / \} - [/ (p(xi/1^a) - 1/N) / \cdot \neg\alpha] = \text{zero or positive; reliable increase in equality is accepted}$$

Contraste de Variación Posible Fiable Empírica Individual, en estudio de sesgo positivo, mediante Segunda Modalidad de Fiabilidad Empírica

$$(p(xi/2^a) - p(xi/1^a)) - [(1 - p(xi/1^a)) \cdot \beta] = \text{zero or positive; reliable increase in positive bias is accepted}$$

Contraste de Variación Posible Fiable Empírica Individual, en estudio de sesgo negativo, mediante Segunda Modalidad de Fiabilidad Empírica

$$(p(xi/1^a) - p(xi/2^a)) - [(p(xi/1^a) - 0) \cdot \beta] = \text{zero or positive; reliable increase in negative bias is accepted}$$

Contrastes de Variación Posible Fiable Empírica Muestral en igualdad de oportunidades, en Primera Modalidad de Fiabilidad Empírica

$$(DM/1^a - DM/2^a) - \{ DM/1^a \cdot \neg\alpha \} = \text{zero or positive; reliable trend toward equality of opportunities is accepted}$$

Contrastes de Variación Posible Fiable Empírica Muestral normal en estudio de sesgo, en Segunda Modalidad de Fiabilidad Empírica

$$(DM/2^a - DM/1^a) - \{ \{ [(1 - 1/N) \cdot 2] : N \} - DM/1^a \} \cdot \beta \} = \text{zero or positive; accepted in bias study}$$

32. Individual projection

At the moment we have what the necessary variation between two points in space should be, it is possible to make projections from the variation between these two points, which is normally called slope in geometry and is the linear estimation of, given a derivative, what the linear prediction is following the same trajectory of the derivative in space.

In the field of Impossible Probability, the difference between an empirical probability and a statistic is the necessary variation for that subject or option to become equal to the statistic. Thus, if the statistic is the Maximum Empirical Probability (Theoretically)

Possible, that is, probability one, in studies of positive bias, then the difference of the unit minus the empirical probability is the variation that the subject or option needs to increase to reach the point of the Maximum Empirical Probability (Theoretically) Possible. In equality studies this necessary variation is the same normal Bias Level. And in studies of negative bias that necessary variation is in itself the empirical probability that must be reduced until reaching zero.

At the moment we understand that the necessary variation is not only a comparative differential between two measurements, but the necessary variation that in space exists between two points for linear estimates of growth or decrease, then is when the concept of individual linear projection arises.

The individual linear projection is nothing more than, given a necessary variation, what is the projection we make for a next measurement of a subject or option in a critical or empirical way. If we understand the starting point for the initial estimation as the first measurement upon which the necessary linear projection for future measurements is established, necessarily the data we are projecting are the projection of data for the second measurement.

The projected probability for the second measurement will be denoted with the symbol: $p(x_i/2^a)'$. The symbol we have adopted for the projection of the second measurement is simply empirical probability of the second measurement “prime”.

In a critical way we can establish what is the sufficiently rational possible variation to accept that a change of positions is sufficiently rational in a second measurement, by comparing the projected empirical probability prime and the real empirical probability obtained when carrying out in practice the second measurement. If the real empirical probability of the second measurement is within the limits of the projected empirical probability prime, the second measurement will be accepted as rational growth or decrease according to the object of study. In this case, it is the scientific policy that establishes the critical reason for the contrast.

At the moment we have the calculation of what part of the necessary variation is acceptable as sufficiently rational in the contrast, equal to the starting point plus the product of the necessary variation by a critical reason, even to that estimation it would

be necessary to add margins of error, understood the margin of error as a margin of oscillation in which the possible future value of subject or option may fit.

The critical reason that multiplies the necessary variation will be denoted with the symbol " $p(xc/a)$ ", and is the critical probability that multiplies the necessary variation indicating what part of the necessary variation we consider sufficiently rational to be added to the starting point, the initial subject or option empirical probability. To the result of the sum the error margins are added; the margin of error in the projection will be symbolised " $p(xc/b)$ ", and that margin of error is the product of a percentage X of error or reliability multiplied by the maximum trend to which the necessary variation may aspire. The way to understand this margin of error is by establishing the lower error margin equal to the difference of the critical probability " $p(xc/b)$ " minus the result of the sum of the empirical probability plus the critically possible variation. And the upper margin equal to the sum of that possible variation plus the critical probability " $p(xc/b)$ ".

Normally, when comparing the data of a second measurement in relation to the prime projection, as long as the real data of the second measurement are within the lower or upper limits, that is, it is equally and simultaneously fulfilled that they are: equal or greater than the lower limit, and equal or less than the upper limit, the data of the second measurement would be accepted rationally. In growth studies it could be the case to accept only and exclusively values equal or greater than the upper limit to increase the requirement. Likewise, in studies of decrease one could accept only data in a second measurement equal or less than the projected lower limit.

In addition to rational critical projections on critically possible variation, empirical projections could be made on empirically possible variations, including in any case empirical lower or upper error limits on the projected empirical values.

Empirically, if we know what the degree of possible empirical sample variation is, we can establish that empirical criterion as a criterion of possible empirical variation, so that the product of the necessary variation by the empirical variation criterion will be equal to the expected possible empirical variation. Those empirical variation criteria will again be the first and second modalities of empirical error or reliability explained previously. Even in the case of calculating the possible variation empirically, to the result of the sum of the

empirical probability plus the possible empirical variation, the lower and upper contrast margins must be added by the sum or subtraction of the critical probability “ $p(xc/b)$ ”.

In the case that the projection, rational or empirical, were to contrast the results between the projected probability and the results of a second measurement of different N, then apply the effect of N. The exposition of the formulation including the effect of N is also developed in Introduction to Impossible Probability.

In addition, there is the possibility of making statistical projections for more than one successive measurement, which will be called emesimal measurements. The calculation for each emesimal projection is also developed in Introduction to Impossible Probability, including the possibility of emesimal measurements of different N sample magnitude.

Projected Critical Probability, in equality of opportunities study, to reduce positive bias

$$p(xi/2^a)' = \{ p(xi) - [(p(xi) - 1/N) \cdot p(xc/a)] \} \pm p(xc/b)$$

$$p(xc/a) = Xa : 100$$

Xa = reliability percentage

$$p(xi/2^a)' = \{ p(xi) - [(p(xi) - 1/N) \cdot p(xc/a)] \} + p(xc/b) = \text{upper limit of the error margin}$$

$$p(xi/2^a)' = \{ p(xi) - [(p(xi) - 1/N) \cdot p(xc/a)] \} - p(xc/b) = \text{lower limit of the error margin}$$

$$\pm p(xc/b) = \{ 1/N - \{ p(xi) - [(p(xi) - 1/N) \cdot p(xc/a)] \} \} / (Xb : 100)$$

Xb = error percentage

Projected Critical Probability, in equality of opportunities study, to reduce negative bias

$$p(xi/2^a)' = \{ p(xi) + [(1/N - p(xi)) \cdot p(xc/a)] \} \pm p(xc/b)$$

$$p(xc/a) = Xa : 100$$

Xa = reliability percentage

$$p(xi/2^a)' = \{ p(xi) + [(1/N - p(xi)) \cdot p(xc/a)] \} + p(xc/b) = \text{upper limit of the error margin}$$

$$p(xi/2^a)' = \{ p(xi) + [(1/N - p(xi)) \cdot p(xc/a)] \} - p(xc/b) = \text{lower limit of the error margin}$$

$$\pm p(xc/b) = \{ 1/N - \{ p(xi) + [(1/N - p(xi)) \cdot p(xc/a)] \} \} \cdot (Xb : 100)$$

Xb = error percentage

Projected Critical Probability, in bias study, to increase empirical probability to Maximum Empirical Probability Possible

$$\begin{aligned}
 p(x_i/2^a)' &= \{ p(x_i) + [(1 - p(x_i)) \cdot p(x_c/a)] \} \pm p(x_c/b) \\
 p(x_c/a) &= X_a : 100 \\
 X_a &= \text{reliability percentage} \\
 p(x_i/2^a)' &= \{ p(x_i) + [(1 - p(x_i)) \cdot p(x_c/a)] \} + p(x_c/b) = \text{upper limit of the error margin} \\
 p(x_i/2^a)' &= \{ p(x_i) + [(1 - p(x_i)) \cdot p(x_c/a)] \} - p(x_c/b) = \text{lower limit of the error margin} \\
 \pm p(x_c/b) &= \{ 1 - \{ p(x_i) + [(1 - p(x_i)) \cdot p(x_c/a)] \} \} \cdot (X_b : 100) \\
 X_b &= \text{error percentage}
 \end{aligned}$$

Projected Critical Probability, in bias study, to reduce empirical probability to Minimum Empirical Probability Possible

$$\begin{aligned}
 p(x_i/2^a)' &= \{ p(x_i) - [p(x_i) \cdot p(x_c/a)] \} \pm p(x_c/b) \\
 p(x_c/a) &= X_a : 100 \\
 X_a &= \text{reliability percentage} \\
 p(x_i/2^a)' &= \{ p(x_i) - [p(x_i) \cdot p(x_c/a)] \} + p(x_c/b) = \text{upper limit of the error margin} \\
 p(x_i/2^a)' &= \{ p(x_i) - [p(x_i) \cdot p(x_c/a)] \} - p(x_c/b) = \text{lower limit of the error margin} \\
 \pm p(x_c/b) &= \{ p(x_i) - [p(x_i) \cdot p(x_c/a)] \} \cdot (X_b : 100) \\
 X_b &= \text{error percentage}
 \end{aligned}$$

Empirical Projected Probability, in equality of opportunities study, to reduce positive bias. First modality of Empirical Reliability

$$\begin{aligned}
 p(x_i/2^a)' &= \{ p(x_i) - [(p(x_i) - 1/N) \cdot \neg \alpha] \} \pm p(x_c/b) \\
 p(x_i/2^a)' &= \{ p(x_i) - [(p(x_i) - 1/N) \cdot \neg \alpha] \} + p(x_c/b) = \text{upper limit of the error margin} \\
 p(x_i/2^a)' &= \{ p(x_i) - [(p(x_i) - 1/N) \cdot \neg \alpha] \} - p(x_c/b) = \text{lower limit of the error margin} \\
 \pm p(x_c/b) &= / \{ 1/N - \{ p(x_i) - [(p(x_i) - 1/N) \cdot \neg \alpha] \} \} / \cdot (X_b : 100) \\
 X_b &= \text{error percentage}
 \end{aligned}$$

Empirical Projected Probability, in equality of opportunities study, to reduce negative bias. First modality of Empirical Reliability

$$\begin{aligned}
 p(x_i/2^a)' &= \{ p(x_i) + [(1/N - p(x_i)) \cdot \neg \alpha] \} \pm p(x_c/b) \\
 p(x_i/2^a)' &= \{ p(x_i) + [(1/N - p(x_i)) \cdot \neg \alpha] \} + p(x_c/b) = \text{upper limit of the error margin} \\
 p(x_i/2^a)' &= \{ p(x_i) + [(1/N - p(x_i)) \cdot \neg \alpha] \} - p(x_c/b) = \text{lower limit of the error margin} \\
 \pm p(x_c/b) &= \{ 1/N - \{ p(x_i) + [(1/N - p(x_i)) \cdot \neg \alpha] \} \} \cdot (X_b : 100) \\
 X_b &= \text{error percentage}
 \end{aligned}$$

Empirical Projected Probability, in bias study, to increase empirical probability to Maximum Empirical Probability Possible. Second modality of Empirical Reliability

$$\begin{aligned}
 p(x_i/2^a)' &= \{ p(x_i) + [(1 - p(x_i)) \cdot \beta] \} \pm p(x_c/b) \\
 p(x_i/2^a)' &= \{ p(x_i) + [(1 - p(x_i)) \cdot \beta] \} + p(x_c/b) = \text{upper limit of the error margin} \\
 p(x_i/2^a)' &= \{ p(x_i) + [(1 - p(x_i)) \cdot \beta] \} - p(x_c/b) = \text{lower limit of the error margin} \\
 \pm p(x_c/b) &= \{ 1 - \{ p(x_i) + [(1 - p(x_i)) \cdot \beta] \} \} \cdot (X_b : 100) \\
 X_b &= \text{error percentage}
 \end{aligned}$$

Empirical Projected Probability, in bias study, to reduce empirical probability to Minimum Empirical Probability Possible. Second modality of Empirical Reliability

$$\begin{aligned}
 p(x_i/2^a)' &= \{ p(x_i) - [p(x_i) \cdot \beta] \} \pm p(x_c/b) \\
 p(x_i/2^a)' &= \{ p(x_i) - [p(x_i) \cdot \beta] \} + p(x_c/b) = \text{upper limit of the error margin} \\
 p(x_i/2^a)' &= \{ p(x_i) - [p(x_i) \cdot \beta] \} - p(x_c/b) = \text{lower limit of the error margin} \\
 \pm p(x_c/b) &= \{ p(x_i) - [p(x_i) \cdot \beta] \} \cdot (X_b : 100) \\
 X_b &= \text{error percentage}
 \end{aligned}$$

33. Muestral projection

Basically, the projection will be carried out over the linear estimation of what the expected data in a second measurement would be according to our object of study. If in the experimental field our object of study is equality of opportunities, what we rationally expect is that the data of the second measurement will be found close to the inversion of N. To make this contrast in the second measurement we make a linear estimation of what the projected empirical probability prima for that second measurement is, which each subject or option should obtain depending on its starting point, its original empirical probability. According to the individual projection of each isolated individual, we make an individual contrast for each subject or option independently.

It may be the case that in a dataset in a second measurement there are subjects or options that indeed fulfil the individual projections, while other subjects or options do not fulfil their individual projections. The way to estimate that, even though there are subjects or options that do not adjust to the expected projection, nevertheless at the sample level it can still be accepted that the sample as a whole, except for counted exceptions, is found within the rational project, is by comparing the dispersion of the

second measurement in relation to a projected dispersion. That projected dispersion is the sample projection.

From knowing what the necessary variation is between the Mean Deviation and the Maximum Theoretical Mean Deviation Possible in growth, or in equality of opportunities understanding the Mean Deviation itself as necessary variation, we can make estimations of rational or empirical possible variation, and over those rational or empirical possible variations we can make sample projections, so that if we have a sample projection for the second measurement, at the moment we have the real data of that second measurement, it is simply to contrast our experimental project with the real data. If the real data are found within the error margins of the project, the second measurement is accepted within the project; otherwise, it is rejected.

The projected second sample measurement will be symbolised $DM(2^a)'$, meaning Mean Deviation second measurement prima. Indicate again that, if one wishes to use Variance or Standard Deviation, its exposition is included in *Introduction to Impossible Probability*, simply taking into account that the real and possible variations will be calculated over the difference of the Variance or Standard Deviation in relation to the Maximum Theoretical Variance Possible and the Maximum Theoretical Standard Deviation Possible, as well as also for the calculation of rational critical values.

Also indicate that in *Introduction to Impossible Probability* the exposition is made including the effect of N in case the second measurement had a different N, as well as the projection of muestral data for emesimal measurements.

Mean Deviation Projected Critically, in a study of bias, for a second measurement

$$DM(2^a)' = \{ DM + \{ [\{ (1 - 1/N) \cdot 2 \} : N \} - DM] \cdot p(xc/a) \} \pm p(xc/b)$$

$$p(xc/a) = Xa : 100$$

$$Xa = \text{percentage of reliability}$$

$$\pm p(xc/b) = \{ \{ [(1 - 1/N) \cdot 2 \} : N \} - \{ DM + \{ [\{ (1 - 1/N) \cdot 2 \} : N \} - DM] \cdot p(xc/a) \} \} \cdot (Xb : 100)$$

$$Xb = \text{percentage of error}$$

Mean Deviation Projected Critically, in a study of equality of opportunities, for a second measurement

$$\begin{aligned} DM(2^a)' &= \{ DM - [DM \cdot p(xc/a)] \} \pm p(xc/b) \\ p(xc/a) &= Xa : 100 \\ Xa &= \text{percentage of reliability} \\ \pm p(xc/b) &= \{ DM - [DM \cdot p(xc/a)] \} \cdot (Xb : 100) \\ Xb &= \text{percentage of error} \end{aligned}$$

Mean Deviation Projected Empirically, in a study of equality of opportunities, for a second measurement. First Modality of Empirical Reliability

$$\begin{aligned} DM(2^a)' &= \{ DM - [DM \cdot \neg\alpha] \} \pm p(xcb) \\ \pm p(xcb) &= \{ DM - [DM \cdot \neg\alpha] \} \cdot (Xb : 100) \\ Xb &= \text{percentage of error} \end{aligned}$$

Mean Deviation Projected Empirically, in a study of bias, for a second measurement. Second Modality of Empirical Reliability

$$\begin{aligned} DM(2^a)' &= \{ DM + \{ [[(1 - 1/N) \cdot 2] : N \} - DM] \cdot \beta \} \} \pm p(xc/b) \\ \pm p(xc/b) &= \{ \{ [(1 - 1/N) \cdot 2] : N \} - \{ DM + \{ [[(1 - 1/N) \cdot 2] : N \} - DM] \cdot \beta \} \} \} \cdot (Xb:100) \\ Xb &= \text{percentage of error} \end{aligned}$$

34. Individual and Sample Forecast

Projections are characterised just as estimations of data for a second measurement or an m-th measurement, having only a first measurement and estimating the necessary variation between the data of the first measurement and the ideal objective. In essence, the necessary variation is the comparison between our initial starting point and where we want to go — whether to increase or decrease the data in order to increase or reduce bias, depending on whether the ideal is equality, positive bias, or negative bias. According to our ideals and having a starting point (first measurement), we construct the projections.

Now, once the projection phase is surpassed and we already have concrete data in a second measurement, independently of the projection — even if the real data in the second measurement do not conform to the projection — we can still make a forecast of

what the evolution of the data will be in successive measurements, even if they do not correspond to our projection.

My project right now is to travel to Edinburgh in a few hours; the flight is scheduled for two in the afternoon, and during the weekend my project is to travel to Inverness and Aberdeen if the weather allows.

All these projects I make in relation to a meteorological forecast according to the meteorological conditions right now, that at this moment suggests that the airplane will take off without problems from Belfast International Airport, and that there will be no blockage on the train lines during the weekend. That is, I have made a project on the basis of current meteorological data.

However, the weather in the north of the United Kingdom is predictable only within margins of error within which, in the North Atlantic, everything can change in a matter of hours. In fact, I have already seen some news anticipating early snowfalls this autumn in the north of the United Kingdom, which will surely affect Scotland.

Regardless of what my project is, the real data in second measurements may or may not conform to my project, and may even invalidate my project. It is possible that this afternoon's flight may be cancelled, or that the weather in Scotland may worsen and make my trip to Inverness and Aberdeen impossible. In any case, although the data of a second meteorological measurement do not correspond to my project, and even if in the extreme case I had to cancel my project, the data we have about the weekend are the data that exist, and any future project I may make — regardless of the fact that the original project has been cancelled — must be made on the basis of the real data available.

This is the difference between a projection and a forecast. Every projection is based on a mental construction of a possible future — what we want or desire. A forecast, on the other hand, is based on reality itself, what there is.

Any contradiction between projection and forecast is the contradiction between our plans and what we actually have. For this reason, the meteorological forecast is not a scientific project; it is the real forecast derived from real data.

At the moment when we have more than one measurement at different times, we can establish forecasts about the future evolution of the data, independently of whether the forecasts match our plans.

Forecasts can be made at the individual level as soon as we have data about a subject or option in more than one measurement at different times, and at the sample level when we have collections of data at different moments. In any case, forecasts must adapt to any sample regardless of its size, including the effect of N, which is included in *Introducción a la Probabilidad Imposible*, and forecasts can also be made for m-th measurements.

The legend that will be used is the following: the forecast empirical probability will be symbolised $p(x_i)''$, meaning “empirical probability two prime”. And the forecast Mean Deviation will be symbolised DM'' , “Mean Deviation two prime”.

Forecast Probability for a third measurement, in a study of growth

$$p(x_i)'' = \{ p(x_i/2^a) + [p(x_i/2^a) - p(x_i/1^a)] \} \pm p(x_c) \\ \pm p(x_c) = \{ 1 - \{ p(x_i/2^a) + [p(x_i/2^a) - p(x_i/1^a)] \} \} \cdot (X : 100) \\ X = \text{percentage of error}$$

Forecast Probability for a third measurement, in a study of decrease

$$p(x_i)'' = \{ p(x_i/2^a) - [p(x_i/1^a) - p(x_i/2^a)] \} \pm p(x_c) \\ \pm p(x_c) = \{ p(x_i/2^a) - [p(x_i/1^a) - p(x_i/2^a)] \} \cdot (X : 100) \\ X = \text{percentage of error}$$

Forecast Mean Deviation for a third measurement, in case of growth of bias between the two measurements

$$DM'' = \{ DM/2^a + (DM/2^a - DM/1^a) \} \pm p(x_c) \\ p(x_c) = \{ \{ [(1 - 1/N) \cdot 2] : N \} - \{ DM/2^a + (DM/2^a - DM/1^a) \} \} \cdot (X : 100) \\ X = \text{percentage of error}$$

Forecast Mean Deviation for a third measurement, in case of growth in equality of opportunities between the two measurements

$$DM'' = \{ DM/2^a - (DM/1^a - DM/2^a) \} \pm p(xc)$$

$$p(xc) = \{ DM/2^a - (DM/1^a - DM/2^a) \} \cdot (X : 100)$$

X = percentage of error

35. Individual inter-measurement studies

At the moment we have surpassed the project phase and have entered the forecast phase, in addition to comparing whether the forecast conforms to the project we can compare the data between the different measurements in time.

As previously mentioned, forecast and project are antagonistic: project is what we desire, forecast is what there really is. In case of contradiction between project and forecast, modifications to the project must be made in order to achieve more favourable future forecasts.

If our project is a vaccine for a certain disease, and at the moment we surpass the project phase, in the second measurement the forecast is adverse, it means that we have to cancel the initial project, which in a constructive sense does not necessarily imply the deactivation of the entire project, but the elimination of the parts of the project that do not work, substituting them for new components that imply an improvement in the efficiency and effectiveness of the project.

If a vaccine does not give the expected results, study the interaction of its components with the virus, and substitute the necessary components in order to obtain a higher performance in the vaccine in future measurements.

In the process of improving a project from a second measurement, it is necessary to take into account the individual evolution of each subject or option. In the case of the vaccine, study how each subject or option to whom the vaccine has been supplied has interacted with the vaccine, trying to find reasons for said individual results in the clinical history of the patient or previous conditions.

In the process of modifying the project, individual inter-measurement comparisons, in addition, can have two formats. Not only comparing the individual data of the same subject or option between two measurements in different moments — intra-individual inter-measurement comparison —. In addition, we can make comparisons between data of different subjects and options in the same measurement or in different moments. For example, patient A has responded to the vaccine with results X, and patient B has responded to the same treatment with results Y. The comparison of data X of A and data Y of B can be done by comparing X and Y, which would be inter-individual studies. Inter-individual studies can be intra-measurement — both patients A and B belong to the same measurement — or inter-measurement, when patient A and patient B belong to different measurements in time. Although patient A and patient B belong to different measurements in time — for example, patient A may have been in 1968 and patient B is in 2025 —, if the medical treatment applied to patient A in 1968 and to patient B in 2025 is the same medical treatment, independently of the 57-year time lapse, insofar as both patients have received the same medical treatment, the comparison between our data even in the 21st century would be of utmost transcendence as evidence of the medical results of our technology. Once again indicate that inter-measurement inter-individual comparisons may be affected by the sample size. For example, if the sample in 1968 had a number of patients N higher than the sample N of 2025, it would be necessary to apply the Effect of N to compare the data between both samples. In any case, given that the Analysis of Impossible Probability attempts to be a summary as concise and clear as possible for all types of readers — not only judges and experts, but any reader interested in our work, including students —, simply say that for a more detailed exposition of the formulations, refer to **section 19 of Introduction to Impossible Probability**. In the following expositions we will limit ourselves to rational critical inter-measurement inter-individual comparisons. In inter-individual studies the legend that will be used is as follows:

$p(x_i/1^a)$ = empirical probability of the subject or option of the first measurement

$p(x_n/2^a)$ = empirical probability of a different subject or option of the second measurement

Critical Inter-measurement Inter-individual Growth

$$p(x_c) - [(1 - p(x_i/1^a)) - (p(x_n/2^a) - p(x_i/1^a))] = \text{zero or positive growth is accepted}$$
$$p(x_c) = (1 - p(x_i/1^a)) \cdot (X : 100)$$

X = percentage of error

Critical Inter-measurement Inter-individual Decrease

$$p(x_c) - p(x_n/2^a) = \text{zero or positive decrease is accepted}$$
$$p(x_c) = p(x_i/1^a) \cdot (X : 100)$$

X = percentage of error

Critical Inter-measurement Inter-individual Equality

$$p(x_c) - / (p(x_n/2^a) - 1/N) / = \text{zero or positive equality is accepted}$$
$$p(x_c) = / (p(x_i/1^a) - 1/N) / \cdot (X : 100)$$

X = percentage of error

36. Reliable Significance in Inter-measurement Individual Study

In the same way that in intra-sample inferential statistics, at the beginning of the sections dedicated to statistical inference on a sample (which is in essence the phase prior to the project phase, a project phase that will be followed by the forecast phase), inferences were made using validity and significance, in the inter-measurement study the statistical tests of section 34 could also be reconverted into inter-measurement statistical tests of statistical significance.

Reliable Significance of Inter-measurement Inter-individual Growth

$$(p(x_n/2^a) - p(x_i/1^a)) - [(1 - p(x_i/1^a)) \cdot p(x_c)] = \text{zero or positive growth is accepted}$$
$$p(x_c) = X : 100$$

X = percentage of reliability

Reliable Significance of Inter-measurement Inter-individual Decrease

$$(p(x_i/1^a) - p(x_n/2^a)) - [p(x_i/1^a) \cdot p(x_c)] = \text{zero or positive decrease is accepted}$$

$$p(x_c) = X : 100$$

X = percentage of reliability

Reliable Significance of Inter-measurement Inter-individual Equality

$$\{ / (p(x_i/1^a) - 1/N) / - / (p(x_n/2^a) - 1/N) / \} - [/ (p(x_i/1^a) - 1/N) / \cdot p(x_c)] = \text{zero or positive equality is accepted}$$

$$p(x_c) = X : 100$$

X = percentage of reliability

37. Critical Inter-measurement Individual Proportion

Normally, differential comparisons tend to be made, although in Impossible Probability, when performing intra-sample inferential statistics, individual or sample, critical proportions were already introduced, which in essence is to use the quotient relation between two factors as a comparison model, instead of the differential comparison.

In the differential comparison what is compared is the magnitude in which one value is higher or lower than another reference value. However, in the quotient comparison the real object of study is to determine the magnitude of proportion of a lower value divided by a higher value, or to determine the value X that multiplied by a lower value is equal to a higher one. For example, four divided by two means that the value X equal to two can be multiplied by the lower value 2 to obtain the higher value 4. The logic of comparison by quotient is in essence the basis of the Level of Similarity, equal to the logarithm base 10 between two values, so that the closer the quotient is to one, the closer the Level of Similarity is to zero, thus a greater approximation to zero implies greater similarity.

In inter-measurement studies we can compare two measurements through the difference of their data, or we can make quotient comparisons, dividing data of one measurement by the data of another measurement, and rationally criticise the result of the quotient depending on the object of study, equality or bias, positive or negative.

Critical Inter-measurement Inter-individual Proportion of Positive Bias

$(p(xn/2^a) : p(xi/1^a)) - p(xc) = \text{zero or positive}$ positive bias growth is accepted

$$p(xc) = (1 : p(xi/1^a)) \cdot (X : 100)$$

X = percentage of reliability

Critical Inter-measurement Inter-individual Proportion of Negative Bias

$p(xc) - (p(xn/2^a) : p(xi/1^a)) = \text{zero or positive}$ negative bias growth is accepted

$$p(xc) = X : 100$$

X = percentage of error

Critical Inter-measurement Inter-individual Proportion of Equality

$p(xc) - [/ (p(xn/2^a) - 1/N) / : / (p(xi/1^a) - 1/N) /] = \text{zero or positive}$ equality is accepted

$$p(xc) = X : 100$$

X = percentage of error

38. Inter-measurement Descriptive Statistics Intra-individual

If we observe carefully the entire process as a whole from the first section of this work, we may appreciate that the first phase of the research begins with descriptive statistics, which is based on empirical induction. In the phase of description of reality we simply adhere to what reality itself shows us; any new description of a given reality, if integrated into the application as a new category, would be considered Artificial Research by Application.

From the induction or application phase we move to the phase of data analysis by deduction or inference, Artificial Research by Deduction, which is based on the establishment of inferential hypotheses from the synthesis of which mathematical algorithm adapts to the data collection.

In essence, descriptive statistics is the application phase; inferential statistics is the phase of data analysis by deduction; and, once we have a series of rational hypotheses, we can elaborate models and projects in the third phase of Artificial Intelligence. In this third phase projects are elaborated according to objectives, and

the projects are contrasted as the data evolve, for the elaboration of real forecasts, as the subjects or options evolve, and to reformulate the projects if there were any contradiction with the forecasts.

At the moment that the contradiction between project and forecast leads to the cancellation of the project, in reality that cancellation means reformulation of the project for obtaining new measurements that may provide more favourable forecasts in the future.

Throughout this entire process —first phase: description and categorisation; second phase: contrast and inference; third phase: project and forecast, and within the third phase, the beginning of the cyclic process of reformulation of the project as many times as necessary for obtaining favourable forecasts— throughout this entire process, in the end we have obtained an “nth” series of measurements.

What at the beginning was only one measurement with the sole purpose of describing a phenomenon, in the end has become a process of data collection at different evolutionary moments that has been able to provide different measurements.

At the moment that, for a subject or option, we have different measurements over time —nth measurements—, then we may include all the measurements of that subject or option in a single sample that permits the development of an Inter-measurement Descriptive Statistical Analysis Intra-individual.

In the Inter-measurement Descriptive Statistics Intra-individual what we are going to do in the Second Method is literally to include the empirical probability of that subject or option in the different measurements, and treat all the empirical probabilities of that subject or option as if they were themselves a sample in itself, without changing anything whatsoever.

The empirical probability that the subject or option had in each measurement is taken intact and is included intact (unless the Effect of N must be applied) in the inter-measurement intra-individual sample, and the inter-measurement intra-individual sample is the collection of each intact empirical probability that the subject or option has obtained in each measurement in which it has been part of the study.

The only reason why an empirical probability is not included intact in the evolutionary sample of that subject or option is because the sample N of that empirical probability

corresponds to a different N sample, hence the Effect of N must be applied. If in the set of empirical probabilities more than one or all measurements had a different N magnitude, identify which sample is considered measurement “m”, and consider all the other empirical probabilities “mi” to be applied to the Effect of N.

From considering all the intact empirical probabilities of a same subject or option as part of the data collection of that subject or option in its temporal evolution, to that intact collection of data statistical study techniques may be applied, beginning with descriptive statistical techniques and, potentially, inferential ones.

Although in the Analysis of *Introduction to Impossible Probability*, given that the object is a brief analysis of the principal contributions of our work, we will limit ourselves to indicating the descriptive formulations, indicating that potentially they would be object of inferential analysis.

For the inter-measurement intra-sample descriptive study of intact data of a same subject or option in an evolutionary time period, the following legend will be used:

M = total number of measurements = { m1, m2, m3 me }

me = nth measurement

p(xi/me) = empirical probability of a subject or option in each nth measurement

From this legend the series of formulations developed in *Introduction to Impossible Probability* is presented, indicating once again that, in the case that an empirical probability belonged to a sample of different N, the Effect of N would have to be applied.

Inter-measurement Intra-individual Arithmetic Mean

$$\sum p(xi/me) : M$$

Inter-measurement Intra-individual Mean Deviation

$$\sum / [p(xi/me) - (\sum p(xi/me) : M)] : M$$

Inter-measurement Intra-individual Variance

$$\sum [p(xi/me) - (\sum p(xi/me) : M)]^2 : M$$

Inter-measurement Intra-individual Standard Deviation

$$\sqrt{\{ \sum [p(xi/me) - (\sum p(xi/me) : M)]^2 : M \}}$$

39. Inter-measurement Sample Criticism

In inter-measurement studies, in addition to individual comparisons, whether intra-individual or inter-individual, intra-sample or inter-sample, intra-measurement and inter-measurement studies work in the same way. It is necessary to contrast sample data, in order to determine that the inferences made at the individual level are also applied at the sample level, or, in the case that at the individual level there were a small group of subjects or options for whom inference is not possible, but there exists a considerable group in N where inference is feasible, that inference may be generalisable at the sample level, as long as at the sample level inference is possible, even when there were a small group of subjects or options that did not satisfy the requirements of statistical inference.

For this reason, in order to validate inference not only at the individual level, but also at the sample level, in *Introduction to Impossible Probability* the inter-measurement sample contrast models are also developed. For which the following legend is used:

DM/1^a = Mean Deviation first measurement

DM/2^a = Mean Deviation second measurement

From this legend, the developments applied to Mean Deviation are presented below, noting as always that they may be adapted to Variance and Standard Deviation, and if necessary to include the Effect of N.

Reliable Critical Proportion of Inter-measurement Sample Growth

$$(DM/2^a : DM/1^a) - p(xc) = \text{zero or positive empirical dispersion growth is accepted}$$
$$p(xc) = \{ \{ [(1 - 1/N) \cdot 2] : N \} : DM/1^a \} \cdot (X : 100)$$

X = percentage of reliability

Inter-measurement Sample Validity of Equality

$$p(xc) - DMoS/2^a = \text{zero or positive equality is accepted}$$
$$p(xc) = DMoS/1^a \cdot (X : 100)$$

X = percentage of error

40. Inter-measurement Omega Individual Studies on a Critical Ratio

The analysis of omega inter-measurement studies will be divided into different sections in order to make their reading easier.

One of the reasons why we decided to write *Analysis of Impossible Probability* is because, when we wrote *Introduction to Impossible Probability*, the project itself was a process of research and self-discovery.

At the moment in which any recording of how *Introduction to Impossible Probability* was written and developed can be publicly shown, it will be perceptible that it was not a linear process; just like every process in my life, it was full of ups and downs, moments of rapid advance and moments of rewriting. In fact, it could be said that it was the moment when I fully formulated the concept of critical probability when all the rational contrast models began to take shape. The idea materialised in my mind one day while reading at the swimming pool in the El Pilar neighbourhood in Madrid. It was at that moment when *Introduction to Impossible Probability* began to take form and meaning.

Likewise, I still remember today how the theory of Maximum Possible Empirical Bias equal to Total Bias divided by two appeared in my mind. I remember perfectly that that day, just when I was in the shower, suddenly the idea burst in with the clarity of a hologram.

Because the very process of constant writing and rewriting of *Introduction to Impossible Probability* was not a linear process, there are extremely dense sections where the explosion of mathematical formulas follows uninterruptedly, while very dense sections

of mathematical philosophy and apparently chaotic organisation appear. However, at the moment I was writing the work, my social activism shifted toward anarchist groups, and I still remember that, while writing the work, on numerous occasions I had certain ideas about the importance of Feyerabend and his work *Against Method*. Although it may seem contradictory, in reality it holds similarities with critical rationalism: what matters is not how scientific ideas arise, but whether they are really valid.

The real point of difference between Feyerabend and critical rationalism lies in the fact that, regardless of whether the origin of scientific hypotheses may be chaotic or erroneous (trial and error demonstrates the chaotic origin of science), although we recognise the chaotic origin of ideas in the mind, we need a scientific method capable of giving rational validation to mental chaos, and this is where we need the rational scientific method as established by Descartes: identifying clear and distinct ideas and being capable of analysing them in their different components, however minimal they may be, in order to reconstruct them afterwards.

Because the task of analysis is, in reality, the identification of elements, while *Introduction to Impossible Probability* was written in 25 very dense sections and filled with successive formulations one after another in many chapters, in the present *Analysis of Impossible Probability* we have adopted the strategy of separating the most representative elements into less dense sections and more focused on the information elements we want to present. This is the reason why we call the present work "analysis": due to the analysis and decomposition of the aspects we consider most noteworthy, presenting them in the most concise way possible.

In the analysis of omega inter-measurement models in the present section we focus exclusively on the contrasts through a critical ratio of possible variation in omega models, and the omega individual projections under critical ratio. Leaving for subsequent independent sections the remaining omega inter-measurement analyses, individual or sample-based, whether of critical probability or empirical criteria.

Just as in normal inter-measurement studies, from the necessary variation we establish what part of the necessary variation is estimated as critically possible to be rationally accepted, this same contrast can be applied to omega models.

And just as, based on necessary variation in normal models, we can elaborate critical projections, which are essentially our project, in omega models we can elaborate projects on the necessary variation in omega models.

Simply recalling that for contrasts between samples of different N the Effect of N must be included in the formulations, and remembering that projections may be made for nth-order measurements.

Contrast of Reliable Theoretical Possible Variation — Individual Omega
difference of the necessary variation in the first measurement minus necessary variation in the second measurement, and the result minus the Reliable Theoretical Possible Individual Omega Variation

$$[/ (p(x_i/1^a) - 1/\Omega) / - / (p(x_i/2^a) - 1/\Omega) /] - [/ (p(x_i/1^a) - 1/\Omega) / \cdot (X : 100)] = \text{zero or positive reliable levelling of the ideals is accepted}$$

X = percentage of reliability

Theoretically Projected Probability on Reliable Theoretical Possible Variation — Individual Omega when in the first measurement the empirical probability of the ideal is higher than the inverse of omega, for a second measurement

$$p(x_i/2^a)' = \{ p(x_i/1^a) - [(p(x_i/1^a) - 1/\Omega) \cdot (X_a : 100)] \} \pm p(x_c)$$

X_a = percentage of reliability

$$\pm p(x_c) = / \{ 1/\Omega - \{ p(x_i/1^a) - [(p(x_i/1^a) - 1/\Omega) \cdot (X_a : 100)] \} \} / \cdot (X_b : 100)$$

X_b = percentage of error

Theoretically Projected Probability on Reliable Theoretical Possible Variation — Individual Omega when in the first measurement the empirical probability of the ideal is lower than the inverse of omega, for a second measurement

$$p(x_i/2^a)' = \{ p(x_i/1^a) + [(1/\Omega - p(x_i/1^a)) \cdot (X_a : 100)] \} \pm p(x_c)$$

X_a = percentage of reliability

$$\pm p(x_c) = \{ 1/\Omega - \{ p(x_i/1^a) + [(1/\Omega - p(x_i/1^a)) \cdot (X_a : 100)] \} \} \cdot (X_b : 100)$$

X_b = percentage of error

41. Omega Inter-measurement Sample Studies on a Critical Ratio

As in normal inter-measurement studies, in addition to individual rational comparisons we need sample comparisons, in order to determine whether, at the sample level, the tendency observed individually is maintained or not. As mentioned on different occasions, there are situations in which not the entire sample follows the ideal behaviour: there may be groups of subjects or options that resist experimental manipulation. Thus, in order to determine that, despite such observations, a proper sample tendency is maintained, a critical sample analysis must be carried out to shed light on statistical decisions.

At the sample level, one must compare that the difference between the empirical sample dispersion and the ideal sample dispersion in the second measurement is lower than that observed in the first measurement, and that difference must be subjected to rational critical contrast.

At the moment we have two different observations of sample tendency, we can also proceed to the critical contrast of the second measurement in relation to the projected dispersion based on our ideal.

Always remember the possibility of adapting the formulas to the Effect of N if necessary, and in projections the possibility of projections for nth-order measurements.

Contrast of Reliable Theoretical Possible Variation — Sample Omega

$$\begin{aligned} & \{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / - \\ & / \{ DM/2^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / - \\ & \{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / \cdot (X : 100) \} = \text{zero or positive reliable results} \\ & \text{are accepted} \\ & X = \text{percentage of reliability} \end{aligned}$$

Theoretically Projected Sample Mean Deviation Omega in second measurement, if the Mean Deviation in the first measurement is higher than Ideal Mean Deviation

$$DM/2^a = \{ DM/1^a - \{ \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \cdot (Xa : 100) \} \} \pm p(xc/b)$$

Xa = percentage of reliability

$$\pm p(xc) = / \{ \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} - \{ DM/1^a - \{ \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \cdot (Xa : 100) \} \} \} \} / \cdot (Xb : 100)$$

Xb = percentage of error

Theoretically Projected Sample Mean Deviation Omega in second measurement, if the Mean Deviation in the first measurement is lower than Ideal Mean Deviation

$$DM/2^a = \{ DM/1^a + \{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / \cdot (Xa : 100) \} \} \pm p(xc/b)$$

Xa = percentage of reliability

$$\pm p(xc) = \{ \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} - \{ DM/1^a + \{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / \cdot (Xa : 100) \} \} \} \cdot (Xb : 100)$$

Xb = percentage of error

42 Third modality of empirical error and empirical reliability in omega models

In the specific case of omega models, empirical error will be equal to the difference between empirical dispersion and ideal dispersion—differential result divided by ideal dispersion. The purpose of this calculation is to determine what the necessary sample variation represents in the omega model in relation to ideal dispersion, so that this proportion is considered error insofar as it represents the proportion by which we have not reached the ideal. That proportion is, in essence, the proportion of empirical sample error in the model, the proportion we must reduce as much as possible to reach the ideal. Therefore, empirical reliability will be equal to one minus that quotient. Empirical reliability in omega models will be called the third modality of empirical reliability and will be symbolised with the representation of not gamma, $\neg\gamma$; thus, the third modality of empirical error is gamma, γ .

Third Modality of Empirical Error, $\neg\gamma$, in DM

$$\neg\gamma = / \{ DM - \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / : \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \}$$

Third Modality of Empirical Reliability

$$\gamma = 1 - \{ / \{ DM - \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / : \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \}$$

43. Intermeditational omega studies at the individual level on empirical reliability

Once we have empirical criteria of error and reliability, they can be used in rational models to establish what part of the necessary variation we consider empirically possible given empirical reliability itself. In this way, empirical reliability can be established as an empirical contrast criterion both in the critique of possible variation and in the formulation of projected probabilities based on the first measurement, using them as axes of contrast in relation to a second measurement.

In order to offer the most significant empirical contrast models, only reliable contrast models and projections based on empirical reliability have been included in this selection. For further development, the reader must refer to the main work, *Introduction to Impossible Probability*.

Empirically Reliable Possible Variation Contrast, Individual Omega

$$[/ (p(\xi/1^a) - 1/\Omega) / - / (p(\xi/2^a) - 1/\Omega) /] - [/ (p(\xi/1^a) - 1/\Omega) / \cdot \gamma] = \text{zero or positive};$$

reliable levelling of the ideals is accepted.

Empirically Projected

Probability on Empirically Reliable Possible Variation, Individual Omega, when in the first measurement the empirical probability of the ideal is higher than the inversion of omega, for a second measurement

$$p(xi/2^a)' = \{ p(xi/1^a) - [(p(xi/1^a) - 1/\Omega) \cdot \gamma] \} \pm p(xc) \\ \pm p(xc) = / \{ 1/\Omega - \{ p(xi/1^a) - [(p(xi/1^a) - 1/\Omega) \cdot \gamma] \} \} / \cdot (Xb : 100) \\ Xb = \text{percentage of error}$$

Empirically Projected Probability on Empirically Reliable Possible Variation, Individual Omega, when in the first measurement the empirical probability of the ideal is lower than the inversion of omega, for a second measurement

$$p(xi/2^a)' = \{ p(xi/1^a) + [(1/\Omega - p(xi/1^a)) \cdot \gamma] \} \pm p(xc) \\ \pm p(xc) = \{ 1/\Omega - \{ p(xi/1^a) + [(1/\Omega - p(xi/1^a)) \cdot \gamma] \} \} \cdot (Xb : 100) \\ Xb = \text{percentage of error}$$

44. Intermeditational omega studies at the sample level on empirical reliability

At the omega sample level in intermeditational studies, the formulation is simply adapted to the inclusion of empirical error and empirical reliability models. In the present selection, only the models based on empirical reliability are included.

Empirically Reliable Possible Variation Contrast, Sample Omega

$$\{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / - \\ / \{ DM/2^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / - \\ \{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / \cdot \gamma \} = \text{zero or positive; reliable values are} \\ \text{accepted}$$

Empirically Projected Deviation in the second measurement, if the Mean Deviation in the first measurement is higher than the Ideal Mean Deviation

$$DM/2^a' = \{ DM/1^a - \{ \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} \cdot \gamma \} \} \pm p(xc/b) \\ \pm p(xc) = / \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} - \\ \{ DM/1^a - \{ \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} \cdot \gamma \} \} / \cdot (Xb : 100) \\ Xb = \text{percentage of error}$$

Empirically Projected Deviation in the second measurement, if the Mean Deviation in the first measurement is lower than the Ideal Mean Deviation

$$\begin{aligned} DM/2^a &= \{ DM/1^a + \{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / \cdot \gamma \} \} \pm p(xc/b) \\ &\pm p(xc) = \{ \{ \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} - \\ &\{ DM/1^a + \{ / \{ DM/1^a - \{ [1/N \cdot (N - \Omega)] \cdot 2 \} : N \} \} / \cdot \gamma \} \} \cdot (Xb : 100) \\ &Xb = \text{percentage of error} \end{aligned}$$

45. Rational intermeditational individual omega criticism

Up to now, rational intermeditational omega studies are based on the comparison of two different measurements; however, what is proposed now is different: it consists of the comparison of: the absolute value of the Level of Bias in relation to the ideal in the second measurement, $/(p(xi/2^a) - 1/\Omega)/$, over a critical probability, $p(xc)$, which, in a certain sense, maintains the same structure as critical validity.

However, the reason why it is not considered simply validity and is instead included in omega intermeditational models is because the critical probability, $p(xc)$, is the percentage of error we are willing to accept over the Level of Bias in relation to the ideal in the first measurement, $p(xc) = [/(p(xi/1^a) - 1/\Omega)] (X/100)$; X = percentage of error.

The logic is quite simple: if, given a first measurement, the Level of Bias between empirical probability and the inversion of omega is, in essence, the error we want to reduce to zero, then in the second measurement this difference should be zero; therefore, the empirical probability should be equal to the inversion of omega, accepting only an error margin equal to the percentage of error we are willing to accept over the difference between the empirical and the ideal in the first measurement.

As in previous occasions, we do not criticise only differential relations, but also ratio relations; therefore, the way to adapt this comparison to ratio criticism is to directly criticise the ratio of the Level of Bias in relation to the ideal of the second measurement over that obtained in the first measurement. Obviously, the more the ratio tends to zero, the greater the approximation to the ideal; therefore, it will only be accepted if it is equal to or less than a critical ratio.

Omega Intermeditational Individual Levelling

$p(xc) - / (p(xn/2^a) - 1/\Omega) / =$ zero or positive, the intermeditational trend toward the ideal is accepted

$$p(xc) = / (p(xi/1^a) - 1/\Omega) / \cdot (X : 100)$$

X = percentage of error

Omega Intermeditational Individual Proportion

$p(xc) - [/ (p(xn/2^a) - 1/\Omega) / : / (p(xi/1^a) - 1/\Omega) /] =$ zero or positive, the trend is accepted

$$p(xc) = X : 100$$

X = percentage of error

46. Rational intermeditational omega muestral criticism

The models of rational intermeditational individual criticism are adapted at the sample level; what would simply be sample validity of the second measurement, at the intermeditational level is adapted to a rational criticism using as critical ratio that critical probability which would have been calculated over the first measurement.

In the same way, we proceed to criticise directly the ratio of the reliability of the second measurement in relation to the reliability of the first measurement. As long as the empirical ratio is equal to or greater than a critical ratio, it is accepted. The empirical ratio here is the third modality of empirical reliability.

$\gamma/2^a$ = Third Modality of Empirical Reliability in the second measurement

$\gamma/1^a$ = Third Modality of Empirical Reliability in the first measurement

Intermeditational Omega Reliability Validity

$(\gamma/2^a - \gamma/1^a) - p(xc) =$ zero or positive, the trend toward levelling the ideals is accepted

$$p(xc) = (1 - \gamma/1^a) \cdot (X : 100)$$

X = percentage of reliability

Intermeditational Omega Critical Proportion of Reliability

$(\gamma/2^a : \gamma/1^a) - p(xc) = \text{zero or positive, the trend toward levelling the ideals is accepted}$

$$p(xc) = (1 : \gamma/1^a) \cdot (X : 100)$$

X = percentage of reliability

47. Critical Omega Muestral Proportion

The establishment of a modality of omega empirical error and omega empirical reliability, the third modality of empirical error or reliability, both in intrameditational and intermeditational studies, intramuestral or intermuestral, can be criticised directly in relation to a critical probability. It is reiterated that it could be valid in any omega study.

Critical Omega Muestral Proportion of Error

$p(xc) - \neg\gamma = \text{zero or positive, the empirical dispersion of the omega model is accepted}$

$$p(xc) = X : 100$$

X = percentage of error

Critical Omega Muestral Proportion of Reliability

$\gamma - p(xc) = \text{zero or positive, the empirical dispersion of the omega model is accepted}$

$$p(xc) = X : 100$$

X = percentage of reliability

48. Individual Impact of the Defect

Up to now, everything we have studied are the implications that the introduction of the Second Method on 16 October 2002 entails within the work *"Intuition and Probability"*, dedicated to the resurrection of Lazarus.

Both the significance of the Second Method and the dedication have transcendence in the technological field in which this theory has been developed. At the moment our technology can be made public, not only will the medical contributions we may offer to

humanity as a whole be important, but even more specifically those related to the field of Global Artificial Intelligence.

The theory of the cerebral hologram, the basis of our technology and a fundamental part of the development of this theory, could be the beginning of cerebral robotisation suggested by José Rodríguez Delgado, which in a first phase could lead to hybridisation between the brain and Artificial Intelligence; in later phases, models of hybridisation between the brain and Artificial Intelligence could advance to the point of testing traditional concepts of brain death. In other words, brain death could be studied, and in a future perhaps still distant, brain death could be overcome. The dedication to the resurrection of Lazarus in the work where the Second Method is born is a very meaningful indication of the ultimate purpose of our research.

But before the Second Method of 16 October 2002, more transcendent is the introduction of the Impact of the Defect in the early morning of 11 September 2001.

Because before the Second Method of 16 October 2002, more transcendent is the introduction of the Impact of the Defect in the early morning of 11 September 2001, since it is the first original mathematical formula of Impossible Probability. In this section we will simply explain how the individual Impact of the Defect is understood, and in the following sections we will advance to the rational critique of the Impact of the Defect. The reason for this analysis is clear: whereas the Impact of the Defect can indeed be attributed to the original formulations of 11 September 2001, the rational critique of the Impact of the Defect must be dated to the advances we made after the summer of 2010.

To understand the individual Impact of the Defect it is necessary to understand that the meaning and purpose of this formula is, given a system, event, or phenomenon, the ranking distribution in ascending order from least severe to greatest severity among the different categories or typologies of defects. The moment the concept of category is mentioned, we are already saying that this ranking is the application of the Impact of the Defect: first stage. Second stage: the calculation of the Impact of the Defect and its rational critique. Third stage: the decisions derived from the rational critique, namely the decisions needed to increase the efficiency and efficacy of the system.

In essence, the three stages of Artificial Intelligence can be summarised as: data, analysis, decisions. These three stages are applicable to any Artificial Intelligence, Global, Specific, or Particular.

Thus, the ranking of positions is the application of the Impact of the Defect; in the ranking, the first position is occupied by the least severe defect, and the last position is occupied by the most severe defect.

N^o = total number of types of defect in the ranking

n^o = n-th position of each type of defect in the ranking

n^{o1} = defect type number one, the least severe

n^{o2} = defect type number two, more severe than the first, less severe than the third

n^{o3} = defect type number three, more severe than the second and less severe than the next one

n^{o+1} = defect type more severe than the previous and less severe than the next

Once the application has been structured in the form of a ranking of defect categories, it can then weight the severity of each defect. Essentially, that weighting is the calculation derived from the quotient of dividing each particular position of the ranking by the total number of categories; in other words, the Weighted Severity (GP, *Gravedad Ponderada*, in Spanish) is equal to the quotient of each ranking position divided by the value of the last position in the ranking.

Weighted Severity = GP = $n^o : N^o$

n^o = n-th position of each type of defect in ascending order of severity

N^o = total number of positions in the ranking

At the moment we have constructed the application with the distribution of defect categories in a ranking, ordered from a first position corresponding to the least severe category to the last position corresponding to the most severe category, and we have calculated the Weighted Severity, GP, then the application can collect data by including in the application the frequency of defects located in each ranking category. As in the Second Method, the raw score or frequency is represented with the symbol “xi”; in Impact

of the Defect the same will be done, so that “xi” is the frequency of defects per category. Thus, the summation of defect frequencies by categories will be symbolised as $\sum x_i$.

x_i = raw score or frequency = frequency of a type of defect in the process or system
 $\sum x_i$ = summation of frequencies

At the moment we know the Weighted Severity (GP) and the frequency of defects per category, then the individual Impact of the Defect per category will be equal to the quotient between the product of individual frequency of that category and its Weighted Severity, and the result of the product must be divided by the summation of all frequencies.

Individual Impact of the Defect = ID = $[(n^o : N^o) \cdot x_i] : \sum x_i$

x_i = frequency or raw score of the type of defect in the process or system

n^o = n-th position of the type of defect in ascending order of severity

N^o = total number of positions in the ranking

$\sum x_i$ = summation of frequencies or raw scores of all defects

The logic of the equation is simple. If the Maximum Weighted Severity corresponds to the last position of the ranking, then the Maximum Weighted Severity is equal to $N^o : N^o = 1$, which leaves intact the frequency of the most severe category; thus, the Maximum Individual Impact of the Defect would occur when, of all possible defects, the maximum possible defect concentrates the total frequencies—in other words, when all errors committed are the maximum error.

If in a system $\sum x_i$ errors of different categories can be committed, as long as not all $\sum x_i$ errors correspond to the maximum error, the system may present errors, but at least there exist errors less severe than others. If in a system all errors committed are maximum errors, evidently that system needs immediate improvement.

If all $\sum x_i$ are concentrated in defect N^o , the last category in the ranking, then conditions of Maximum Absolute Individual Impact of the Defect occur:

Maximum Absolute Individual Impact of the Defect = $[(N^0 : N^0) \cdot \sum x_i] : \sum x_i = 1$

Thus, the Impact of the Defect is a dimension whose spectrum ranges between zero and one. At the moment in which the spectrum of the Impact of the Defect ranges between zero and one, it bears similarities to a probability model, which will allow us to adapt probability concepts to the study of the Impact of the Defect.

Thus, insofar as a probability model, the Impact of the Defect can be calculated as the product of the empirical probability of a category—frequency of the category divided by total frequencies—multiplied by GP, Weighted Severity. In other words: **GP times empirical probability equals the individual Impact of the Defect.**

Individual Impact of the Defect = ID = $p(x_i) \cdot (n^0 : N^0)$

49. Theoretical Individual Statistics in Impact of the Defect

As mentioned in the previous section, an Impact of the Defect equal to one is only possible when absolutely all the errors in a system are attributable to the maximum category of errors. In that case, if all errors without exception are maximum-level errors, it is said that, under conditions of absolute maximum error, the Maximum Absolute Individual Impact of the Defect occurs.

Maximum Absolute Individual Impact of the Defect

$$[(N^0 : N^0) \cdot (x_i = \sum x_i)] : \sum x_i = (\sum x_i : \sum x_i) \cdot (N^0 : N^0) = 1$$

The Maximum Absolute Individual Impact of the Defect occurs when all errors are distributed in the most severe category. This may or may not occur, but it may also be the case that, for each particular and specific category, the maximum Individual Impact of the Defect for that category would be that all errors—the summation of frequencies—be

distributed in that particular and specific category, not necessarily the most severe one; it may be of intermediate or low severity. In such a situation, that particular and specific category, regardless of its severity, would assume all possible errors, which would constitute the maximum error that such a category could assume.

In this more specific case, if all errors are assigned to a single category, whatever it may be, the frequency of that category would be equal to the total frequency; thus its empirical probability would be one. When multiplied by GP, Weighted Severity, the result would be that GP is the Maximum Individual Impact of the Defect for that individual category.

Maximum Individual Impact of the Defect = GP

$$[(n^0 : N^0) \cdot (x_i = \sum x_i)] : \sum x_i = (\sum x_i : \sum x_i) \cdot (n^0 : N^0) = n^0 : N^0 = GP = \text{Weighted Severity}$$

In the case of the Maximum Absolute Individual Impact of the Defect, this statistic is one; in reality, the Maximum Absolute Individual Impact of the Defect is equivalent to the Maximum Empirical Probability (theoretically) Possible, one, 1. Thus, the Minimum Possible Impact of the Defect can only be zero, just as the Minimum Empirical Probability (theoretically) Possible is zero, under the condition that no error has been generated in the system and therefore the frequency of defects is zero.

Minimum Possible Impact of the Defect

$$[(n^0 : N^0) \cdot (x_i = 0)] : (\sum x_i = 0) = (0 : 0) \cdot (n^0 : N^0) = 0$$

In other words, the theory of the Impact of the Defect, generated in the early morning of 11 September 2001, and the Second Method, from 14 October 2002, are both *compatible* and are, in essence, probability models applied to different situations: in one case, to the improvement of the efficiency and efficacy of systems; in the other, to the study of phenomena and events.

50. Critique of the Individual Impact of the Defect

At the moment we know the Maximum Individual Impact of the Defect, equal to GP, Weighted Severity, and given that GP is the maximum possible defect for each individual category regardless of its position in the defect ranking, we can then rationally critique each Individual Impact based on its relative frequency.

If the Individual Impact of the Defect —*empirical probability times GP*— is equal to or less than the critical probability, the percentage of error we are willing to tolerate over GP, then we accept the Impact of the Defect within the permitted margin of error. In other words, it is an error margin that can still be accepted. That is, if it is not accepted, it is because the error exceeds acceptable margins, and that category should be reviewed.

In every system it is normal for defects to occur; the concern lies in determining when a defect is of such magnitude that it requires immediate intervention. The way to make this decision is through the Critique of the Individual Impact of the Defect, in order to assess whether a chain of defects is within acceptable limits or surpasses those limits, requiring an immediate response.

If we have a patient A and a patient B, and the symptoms are organised in a ranking, and a new medical treatment —experimental variable— is applied, if the symptoms of patient A are within permissible limits and under control thanks to the experimental manipulation, then the experimental manipulation may be maintained as long as the symptoms remain acceptable and under experimental control. If patient B, under the same treatment, were to develop symptoms above acceptable limits, the experimental treatment would have to be stopped and another type of medical intervention applied.

Critique of the Individual Impact of the Defect

$$\begin{aligned} p(xc) - \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} &= \text{zero or positive} \rightarrow \text{accepted} \\ p(xc) &= (n^0 : N^0) \cdot (X : 100) \\ X &= \text{percentage of error} \end{aligned}$$

51. Descriptive Statistics of the Defect Impact

In the Second Method, we have the possibility of producing statistics relative to different statistical, empirical, or theoretical criteria beyond the statistical norm, the arithmetic mean. In Defect Impact, the same could likely be done; however, we have normally focused on constructing a statistic for the Defect Impact similar to the standard one, but adapted to the Defect Impact.

Given that the Individual Defect Impact bears similarities to empirical probability—inasmuch as both oscillate within a spectrum between zero and one—they may have analytical methods that are in some sense similar, for example in critical study and decision-making. Nevertheless, the main difference lies in the fact that, in Defect Impact, the probability is weighted: the empirical probability is multiplied by the Weighted Severity, which means that its behaviour is mediated by the mathematical weighting introduced by GP.

Under these conditions, what we did between 2009 and 2011 was limit ourselves to a normal descriptive statistic of the Defect Impact, although we have always noted that, potentially, Impossible Probability opens many doors to statistical and probabilistic research. The descriptive statistics of the Defect Impact is one possible statistic, but potentially relative statistics of the Defect Impact could also be developed.

What is presented in *Introducción a la Probabilidad Imposible* as the descriptive statistics of the Individual Defect Impact is based on the computation of its average, and, on that average, the calculation of the statistical dispersion relative to the average of individual Defect Impacts. What is presented in *Introducción a la Probabilidad Imposible* as descriptive statistics of the Individual Defect Impact is thus based on the estimation of its average and the statistical dispersion relative to that average.

Average Defect Impact = IDP = $\sum \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} : N^0$

Mean Deviation of the Defect Impact

$$DM^0 = \sum / \{ \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} - \{ \sum \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} : N^0 \} \} / : N^0$$

Variance of the Defect Impact

$$S^{0^2} = \sum \{ \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} - \{ \sum \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} : N^0 \} \}^2 : N^0$$

Standard Deviation of the Defect Impact

$$S^0 = \sqrt{\sum \{ \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} - \{ \sum \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} : N^0 \} \}^2 : N^0}$$

52. Theoretical Sample Statistic of the Defect Impact, 1/N⁰

At the moment we have a descriptive statistic of the Defect Impact, we have automatically generated empirical sample statistics; however, in the critical theory developed in Impossible Probability since the summer of 2010, it is necessary to have theoretical statistics in order to proceed with rational criticism.

Normally, rational criticism is based on the calculation of the percentage of error or reliability we are willing to accept, for which we multiply such critical percentage of error or reliability by statistics of maximum tendency.

Thus, to carry out rational sample criticism of empirical sample statistics, we need theoretical statistics of maximum tendency.

The most important theoretical statistic of maximum tendency will be the inversion of the value N⁰, which has two functions: N⁰ is at the same time the maximum defect category in the defect ranking—because it is the last category—but also, insofar as it is the last category, N⁰ is equal to the value N of categories in the ranking. Therefore, in essence, the

inversion of N^0 will serve multiple functions, since $1/N^0$ in a certain way assumes functions of $1/N$, but in addition $1/N^0$ includes new functions.

If the Maximum Absolute Individual Defect is equal to one, then under conditions of Maximum Absolute Individual Defect the average of Defect Impacts will be equal to dividing the Maximum Absolute Individual Defect by N^0 , which is the same as saying that, under conditions of Maximum Absolute Individual Defect, the average is equal to the inversion of the total number of categories, $1/N^0$; thus, the inversion of N^0 is equal to the Average of Maximum Absolute Individual Defect.

$$1/N^0 = \text{Average of Maximum Absolute Individual Defect}$$

53. Sample Critique of the Defect Impact

From the very moment we have a theoretical sample statistic, we have a statistic of maximum tendency on which to calculate the critical percentage X of error determined by scientific policy.

Thus, the critique of the sample is summarised in verifying whether the Average Defect Impact is equal to or less than the established critical ratio.

Critique of the Average Defect Impact

$$p(xc) - \{ \sum \{ [(n^0 : N^0) \cdot x_i] : \sum x_i \} : N^0 \} = \text{zero or positive, it is accepted}$$

$$p(xc) = 1/N^0 \cdot (X : 100)$$

$$X = \text{percentage of error}$$

$$1/N^0 = \text{Average of Maximum Absolute Individual Defect}$$

$$\text{Maximum Absolute Individual Defect} = 1$$

54. Rational Intermeditational Critique, Individual and Sample

Once the intrameditional studies, individual and sample, have been carried out, decisions can be made on the data to increase the system's efficiency or effectiveness. At the moment these measures are implemented, it is time to take new measurements to verify that the data in the second measurement, after having taken the pertinent measures for the elimination of defects, have had the desired effect on the system as a whole.

Thus, the purpose of the intermeditational study of defect is to verify whether the results obtained from the decisions adopted correspond to the aim of our work: the improvement of our efficiency.

For this reason, individual and sample defects in the second measurement, after the elimination of defects in the system, should present more positive data in terms of significant reduction of defects.

The intermeditational critical study of defects must be carried out both at the individual level and at the sample level. The way this intermeditational individual study will be done is by analysing whether the reduction of defects in any subsequent m -th measurement, in relation to the first measurement, falls within the acceptable margin of error. If not, the defect-elimination policy must be increased until, in later m -th measurements, an elimination of error is achieved consistent with our expectations for the system's efficiency.

At the average level, the same comparison is made by analysing whether the overall reduction of average defects has produced the expected results. Otherwise, at the sample level, general decisions must be taken to eliminate whatever defects may reduce our efficiency.

Rational Intermeditational Critique of Individual Defect Impact

$\{ \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} - \{ [(n^0 : N^0) \cdot xi/me] : \sum xi/me \} \} - p(xc) = \text{zero or positive,}$
reduction of Defect Impact is accepted

$$p(xc) = \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} \cdot (X : 100)$$

X = percentage of reliability

$xi/1^a$ = direct score of Individual Defect Impact in the first measurement

$\sum xi/1^a$ = sum of direct scores in the first measurement

xi/me = direct score of Individual Defect Impact in the “me” measurement

$\sum xi/me$ = sum of direct scores in the “me” measurement

Rational Intermeditational Critique of Individual Defect Impact in “me” measurement=

Rational Intermeditational Critique of Average Defect Impact

$\{ \{ \sum \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} : N^0 \} - \{ \sum \{ [(n^0 : N^0) \cdot xi/me] : \sum xi/me \} : N^0 \} \} - p(xc) =$
zero or positive, reduction of the Average Defect Impact is accepted

$$p(xc) = \{ \sum \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} : N^0 \} \cdot (X : 100)$$

X = percentage of reliability

55. Effective Distribution: Individual Effectiveness

In section 22 of *Introduction to Impossible Probability* we studied the Effective Distribution; in the present analysis it corresponds to section 55. As can be seen, the existence of a larger number of sections is due to the structure of this edition itself, whose purpose is the analysis of the main contributions of *Introduction* in a clearer and more concise format.

In reality, the Effective Distribution follows exactly the same structure as Defect Impact, but adapted to the study of efficacy and efficiency. Essentially, the formulas and equations are the same; what changes is the content.

The reason for this phenomenon is obvious: mathematical formulas, as formulas, are formal structures; the content of the formula is what gives them meaning in one direction

or another. If the content of a formula is defects, what that formal structure analyses are defects; but if the content of that same formula is elements of efficacy and efficiency, then that formula analyses efficacy and efficiency.

Mathematics, as a formal science, will have meaning according to the contents for which they are used. If an algorithm such as addition or subtraction is used to add or subtract candies when teaching mathematics to school children, the interpretation of that mathematical calculation conforms to the psychological structure of those children. If the algorithms of addition or subtraction are used to calculate a worker's salary or the taxes paid by a business owner, the interpretation radically changes.

At the moment in which the formal structures —the formulas— are used to study contents referring to different objects, interpretation automatically, although not really hermeneutic, but still interpretation, becomes adapted to what we are studying.

If a formula is used to study defects, the objective is to eliminate defects, and a greater number of defects will be interpreted as a negative element for our efficiency. On the contrary, if the content consists of elements that stimulate the system's efficacy or efficiency, the objective is to increase the mathematical values that represent this increase in efficacy and efficiency in our system.

At the individual level, the ranking structure of the categories of efficacy and efficiency is organised in ascending order, such that the first category "n°1" corresponds to the category that contributes least to the system's efficacy or efficiency, and the last category, "N°", corresponds to the category of maximum efficacy or efficiency. In this way, the distribution of the ranking is identical but with a different meaning.

N° = total number of n-th categories in the ranking

n° = n-th position of the category in the ranking

$n^{\circ}1$ = category number one of the ranking, the least important category of efficacy

$n^{\circ}2$ = category number two, more important than the first and less than the third

$n^{\circ}3$ = category number three, more important than the second and less important than the fourth

$n^{\circ}+1$ = category more important than the previous and less important than the next

Thus, the quotient n^o / N^o will now be Weighted Importance, IP (*Importancia Ponderada*, in Spanish), such that the category of maximum importance is category N^o ; therefore, the maximum weighting corresponds to N^o , since $N^o / N^o = 1$.

$$\text{Weighted Importance} = IP = n^o / N^o$$

Once we know the Weighted Importance = IP, then we can determine the Individual Effectiveness, EI (*Efectividad Individual*, in Spanish), which will be equal to the product of IP by the empirical probability; in other words, Weighted Importance times direct score or frequency, and the result divided by the sum of direct scores or frequencies.

$$\text{Individual Effectiveness} = EI = \{ [(n^o : N^o) \cdot x_i] : \sum x_i \}$$

$$\text{Individual Effectiveness} = EI = p(x_i) \cdot (n^o : N^o)$$

56. Descriptive Statistics of Individual Effectiveness

The descriptive statistics of Individual Effectiveness will be exactly the same as for Defect Impact; the only thing that changes is the content. Now we are studying effectiveness, therefore the interpretation of the results will be different: the greater the magnitude of the results, the greater the efficiency.

$$\text{Average Effectiveness} = \text{Average Effectiveness} = \sum \{ [(n^o : N^o) \cdot x_i] : \sum x_i \} : N^o$$

Mean Deviation of the Effective Distribution

$$DM^o = \sum / \{ \{ [(n^o : N^o) \cdot x_i] : \sum x_i \} - \{ \sum \{ [(n^o : N^o) \cdot x_i] : \sum x_i \} : N^o \} \} : N^o$$

Variance of the Effective Distribution

$$S^o^2 = \sum \{ \{ [(n^o : N^o) \cdot x_i] : \sum x_i \} - \{ \sum \{ [(n^o : N^o) \cdot x_i] : \sum x_i \} : N^o \} \}^2 : N^o$$

Standard Deviation of the Effective Distribution

$$S^o = \sqrt{\sum \{ \{ [(n^o : N^o) \cdot x_i] : \sum x_i \} - \{ \sum \{ [(n^o : N^o) \cdot x_i] : \sum x_i \} : N^o \} \}^2 : N^o}$$

57. Theoretical Statistics of Individual Effectiveness

In the same way, the theoretical statistics of Effectiveness maintain the formal structure of defect, simply changing the interpretation; now the ideal is to tend toward the maximum possible effectiveness.

Maximum Individual Effectiveness is IP, Weighted Importance, and Maximum Absolute Individual Effectiveness will be equal to one; therefore, its average will be the inversion of N^0 , that is, $1 / N^0$. Finally, Minimum Individual Effectiveness can only be zero.

The concept of Maximum Absolute Mean Deviation of Effectiveness is also introduced, which is essentially the Maximum Theoretical Mean Deviation Possible applied to Individual Effectiveness, given that N^0 represents ranking, and N^0 is not the same as N .

When ranking is studied, it will be seen how the N^0 positions can decrease and change quickly, which is why this formulation receives a different name, to distinguish it from the normal theoretical maximum dispersion.

$$\text{Maximum Individual Effectiveness} = [(n^0 : N^0) \cdot (xi = \sum xi)] : \sum xi = n^0 : N^0$$

$$\text{Minimum Individual Effectiveness} = 0$$

Maximum Absolute Individual Effectiveness

$$[((n^0 = N^0) : N^0) \cdot (xi = \sum xi)] : \sum xi = (n^0 = N^0) : N^0 = 1$$

Average Maximum Absolute Individual Effectiveness

$$\{ [((n^0 = N^0) : N^0) \cdot (xi = \sum xi)] : (\sum xi = xi) \} : N^0 = 1/N^0$$

Maximum Absolute Mean Deviation of Effectiveness

$$\{ (1 - 1/N^0) + \{ 1/N^0 \cdot (N^0 - 1) \} \} : N^0$$

58. Intermeditational Rational Critique of Individual Effectiveness

In essence, the formulations of Effective Distribution are exactly the same as those of Impact of Defect; in reality, only the content changes—the form is the same. The dialectic between form and content operates in such a way that the same form can offer different interpretations depending on the content. Thus, if the content is defect, the task consists of eliminating the defect within the system. If the content is effectiveness and efficiency, the task consists of maximizing, within our field of operation, our effectiveness and our efficiency.

Basically, the contrast models in Impact of Defect can be applied to Effective Distribution simply by changing the interpretation: now the task consists of increasing our effectiveness in the field of operation.

Among the expressions adaptable to Effective Distribution, the Intermeditational Rational Critique of Individual Effectiveness was selected, which essentially critiques whether the increase in our effectiveness and efficiency in an m-th measurement is sufficiently superior to be considered rational.

$$\begin{aligned} & \{ \{ [(n^0 : N^0) \cdot xi/me] : \sum xi/me \} - \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} \} - p(xc) = \text{zero or positive} \\ & \quad \text{effectiveness increase is accepted} \\ & p(xc) = \{ (n^0 : N^0) - \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} \} \cdot (X : 100) \\ & \quad X = \text{percentage of reliability} \end{aligned}$$

59. Sample Critique in Effective Distribution

At the level of sample critique in Effective Distribution, the object of study is to verify whether, in sample terms, a rational increase in our effectiveness and efficiency has been observed. Thus, in average terms, the average Effectiveness in the sample, in m-th measurements, should be higher than the average Effectiveness in a first measurement.

In the same way, just as in the study of defect, as defects tend to zero the dispersion also tends to zero, in Effective Distribution we will understand that the greater the dispersion,

the greater the increase in effectiveness. What we seek is that all effectiveness elements increase their frequency to the maximum; therefore, if all subjects or options increase their frequency, the dispersion of the sample will increase. Thus, the greater the dispersion of effectiveness, the greater the expansion of effectiveness.

Intermeditational Rational Critique of Average Effectiveness

$$\{ \{ \sum \{ [(n^0 : N^0) \cdot xi/me] : \sum xi/me \} : N^0 \} - \{ \sum \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} : N^0 \} \} - p(xc) =$$

zero or positive average Effectiveness increase is accepted

$$p(xc) = \{ 1/N^0 - \{ \sum \{ [(n^0 : N^0) \cdot xi/1^a] : \sum xi/1^a \} : N^0 \} \} \cdot (X : 100)$$

X = percentage of reliability

Sample Critical Level of Bias in Effectiveness

$$DM^0 - p(xc) = \text{zero or positive increase in bias in the distribution of frequency among effectiveness categories is accepted}$$

$$p(xc) = \{ \{ (1 - 1/N^0) + \{ 1/N^0 \cdot (N^0 - 1) \} \} : N^0 \} \cdot (X : 100)$$

X = percentage of reliability

60. The Statistical Ranking of Empirical Probabilities

In section 18 of *Introduction to Impossible Probability* the studies in ranking are analysed, which, in this more detailed analysis, will be developed in section 60. In general terms, ranking studies in Impossible Probability can be described as a synthesis of Effective Distribution with the Second Method.

Chronologically, the Defect Impact was first designed in the early morning of 11 September 2001; the Second Method, on 16 October 2002; and it was around the summer of 2003 when the concept of Effective Distribution first emerged in the form of Hierarchical Organization.

The idea of Hierarchical Organization consisted simply in adapting the Defect Impact to any ranking study in which the interpretation of the content was positive. That is, the structure was designed so that the Hierarchical Organization could be used in the analysis of any hierarchical system, or any hierarchical organization, or any organic

hierarchical model, in which the hierarchy was ordered in ranking. Hence its name: Hierarchical Organization.

The Hierarchical Organization, in general, is transformed into effectiveness and efficiency in *Introduction to Impossible Probability*, insofar as it is understood that the main quality of any Hierarchical Organization must indeed be effectiveness and efficiency. In essence, our system is a system of effectiveness and efficiency.

Now, at the moment when in *Introduction to Impossible Probability* the synthesis of Hierarchical Organization and empirical probability surpasses the initial limits, and it is intended to study subjects or options as elements ordered in ranking according to their empirical probabilities, in the same way that in a game we can determine which players have finished in first place, second place, third place, etc., then we may observe that, if as an effect of our manipulation in reality the subjects or options tend to identical empirical probabilities, the number of ranking positions N^0 will automatically tend to be drastically reduced, given that subjects or options with the same empirical probability will share the same position in the ranking. For this reason, the ranking study surpasses the initial limits and becomes an independent study in its own right, different from defect or effectiveness.

One of the fundamental differences between Ranking versus Defect Impact and Effective Distribution is that while in Defect Impact and Effective Distribution the name of the studies changes as the content changes, in the ranking studies of section 18 of *Introduction to Impossible Probability* this phenomenon does not occur; the ranking may preserve the same name regardless of the content to which it is applied, and ranking even has the possibility of being adapted to studies where the ranking is ordered in increasing precedence and decreasing precedence.

It is said that the precedence ordering is increasing when the first position of the ranking corresponds to the position of lesser value, and from the first position the positions are ordered from lower to higher value. And it will be said that it is decreasing when it is ordered inversely, so that the first position is for the highest value, and the positions are ordered in decreasing form from the first position of highest value to the last of lower value.

Obviously, the ranking can be applied to ordering different types of values, but in the field of the Second Method it will be applied to ordering empirical probabilities. Thus, in increasing order the ranking orders empirical probabilities from the first position for the lower empirical probability, the minimum empirical probability in the sample, to the last position, which will be the empirical probability of greater value, the maximum empirical probability in the sample. When two or more subjects or options share the same empirical probability, the ranking position corresponding to that empirical probability will be a shared ranking position for as many subjects or options as share the same empirical probability value.

Thus, in increasing order, the ranking is organized as follows:

Increasing precedence ordering:

n^o = n-th position in the ranking

$n^o1 = p(x1) = p(x_{i-}) \rightarrow$ empirical probability of the number one position in ranking

$n^o2 = p(x2) > p(x1) \rightarrow$ empirical probability of the number two position in ranking

$n^o3 = p(x3) > p(x2) \rightarrow$ empirical probability of the number three position in ranking

$n^o = p(xn) > p(x(n-1)) \rightarrow$ in increasing order each one higher than the previous

And in decreasing order it will be organized as follows:

Decreasing precedence ordering:

n^o = n-th position in the ranking

$n^o1 = p(x1) = p(x_{i+}) \rightarrow$ empirical probability of the number one position

$n^o2 = p(x2) < p(x1) \rightarrow$ empirical probability of the number two position

$n^o3 = p(x3) < p(x2) \rightarrow$ empirical probability of the number three position

$n^o = p(xn) < p(x(n-1)) \rightarrow$ in decreasing order each one lower than the previous

The total number of positions in the ranking will be symbolized by " N^o ", " N^o = total number of positions in the ranking", and it may be the case in ranking studies that while in a first measurement the value N^o is equal to the value N , as the measurements proceed, by the effect of the experimental manipulation of reality the volume of subjects or options with identical empirical probabilities increases, which will increase the number of shared positions in the ranking, positions shared by as many subjects or options as present the same data in the measurements. Thus, there comes a moment

where the number of positions in the ranking, N^0 , is less than the number of subjects or options N in the sample. Because the mission of N^0 is not to order subjects or options; the mission of N^0 is to order empirical probabilities, which will tend to equalize as a direct effect of the experimental manipulation in reality.

In studies of equality, as the experimental manipulation achieves an increase in equality, a greater number of subjects or options will have empirical probability equal to the inversion of N ; thus all subjects with empirical probability equal to inversion of N will share in the ranking the same shared ranking position; thus N^0 will be less than N . If it were the ideal case that all subjects or options have the same empirical probability equal to inversion of N , then the ranking N^0 would be composed of a single shared position, the only one existing, shared by the entire sample, inversion of N ; thus it would be the case that N^0 equals one, $N^0 = 1$.

In normal positive bias studies, as the ideal subject or option whose empirical probability must tend to one approaches Maximum (Theoretical) Empirical Probability Possible, 1, then all other subjects or options, $N - 1$, will tend towards empirical probability zero; thus the ranking will tend to $N^0 = 2$, one position for all probabilities equal to zero and another position for the single subject or option different from zero.

In studies of samples of zeros, for example, reducing all the symptoms of a disease to zero in a sample of patients, the objective is that the ranking be equal to $N^0 = 1$, a ranking of a single position shared by the entire sample, as the entire sample is zero.

In omega studies, the ideal is that all omega subjects or options be inversion of omega and all non-omega subjects or options be zero; thus the ranking should be $N^0 = 2$, one category shared by all omega subjects equal to inversion omega, and another category for all non-omega subjects or options equal to zero.

Thus, in reality any study of experimental manipulation of reality applying the ranking study will aim to study the reduction of N^0 as our research achieves the objectives established from the very beginning, in the first measurement.

For this reason, even if the changes in the first measurement may be small, as the measurements are repeated, the sample N^0 must be reduced until the desired number of positions is achieved.

The reduction of positions in the ranking in m-th measurements is, in essence, a process of increasing homogeneity in the sample; that is, homogeneity in the sample increases as N^0 tends to one, thus it is inverse to the sample N , so the division of N^0 by N will be understood as index of heterogeneity, which may range between zero and one, understanding zero as zero heterogeneity, understanding one as maximum heterogeneity. This index of heterogeneity is called probability of heterogeneity or homogeneity error in the ranking.

$N^0 : N$ = probability of heterogeneity or homogeneity error in the ranking

If we have an index of heterogeneity, in this case called probability of heterogeneity or homogeneity error, then we can establish theoretical statistics of maxima and minima.

When it is equal to one, conditions of maximum probability of error in homogeneity in the ranking occur. When it is equal to zero, conditions of minimum error occur.

The reason why heterogeneity is considered error is because of the very nature of this type of statistical test: the objective of the experimental manipulation is to reduce the ranking N^0 as much as possible. In models of equality of opportunities and samples of zero the ranking N^0 must consist of a single ranking position, $N^0 = 1$. And in normal positive and omega studies a maximum of two ranking positions, $N^0 = 2$. In ranking studies, heterogeneity is interpreted as error, and homogeneity is interpreted as reliability, so if heterogeneity is error, then reliability would be equal to one minus the probability of error.

$1 - (N^0 : N)$ = probability of reliability in the ranking.

61. Theoretical Statistics of Homogeneity in the Ranking

If we can calculate the probability of heterogeneity or homogeneity error, and it is also the case that, depending on the type of study, we can identify different minimal tendencies, then we can establish the theoretical statistics of heterogeneity or homogeneity error.

Insofar as in equality of opportunities and zero-samples the ideal is the reduction of the number of positions to only one position, then the Minimum Possible Homogeneity Error in models of equality and zero-samples is equal to dividing one by the sample N , $1/N$; thus the inversion of N acquires one more function: inversion of N , $1/N$, equal to Minimum Possible Homogeneity Error in equality and zero-samples, in the ranking.

Then, in positive-bias models and omega models, insofar as the ideal is the reduction of ranking positions to only two positions in the ranking, then the Minimum Possible Homogeneity Error of ranking in positive bias and omega models is equal to dividing two by N , $2/N$.

$1/N$ = Minimum Possible Homogeneity Error of ranking in equality and zero-samples

$2/N$ = Minimum Possible Homogeneity Error of ranking in positive-bias models and omega models

62. Rational Critique of Ranking Homogeneity

At the moment we have theoretical statistics, we can begin to establish critical contrast ratios whose purpose is the verification of empirical propositions. In *Introduction to Impossible Probability*, two intrameasurement contrast models of homogeneity were designed: the Validity of Ranking for studies of equality and zero-samples, and the Critique of Ranking Homogeneity in positive bias and omega.

The reason for the distinction between these two contrast models in ranking studies is directly due to the fact that the theoretical statistics on which the critical reference percentage is obtained are clearly different.

Insofar as in studies of equality of opportunities and zero-samples the tendency is that the number of positions in the ranking is reduced to one, then the theoretical statistic of homogeneity error of reference is the inversion of N, $1/N$. In contrast, in both positive-bias studies and omega models, the ideal is the reduction of positions in the ranking to only two positions; thus the ideal homogeneity error is two over N, $2/N$. At the moment the theoretical statistics used for determining the critical ratio are different, two formulations adapted to the singularities of each contrast model are necessary.

In both cases, although the name of the equation is different, the objective remains the same: to determine to what extent the empirical value of empirical homogeneity in the sample — the quotient of positions in the ranking divided by N — is equal to or less than the critical percentage that scientific policy selects over the Minimum Homogeneity Error in Ranking, whether in equality of opportunities or zero-samples in Validity of Ranking, or in positive bias and omega models in Critique of Ranking Homogeneity.

Validity of Ranking Homogeneity in equality or tendency to zero-samples

$$p(xc) - [(N^0 : N) - 1/N] = \text{zero or positive} \rightarrow \text{accepted}$$

$$p(xc) = 1/N \cdot (X : 100)$$

X = percentage of error

Critique of Ranking Homogeneity in positive bias and omega

$$p(xc) - [(N^0 : N) - 2/N] = \text{zero or positive} \rightarrow \text{accepted}$$

$$p(xc) = 2/N \cdot (X : 100)$$

X = percentage of error

63. Effect N^o of the Ranking

The ranking, far from being a static and constant phenomenon, is a dynamic and changing event, such that the magnitude of the number of positions in the ranking may vary very rapidly at the moment when shared ranking positions begin to increase due to more than one subject or option having the same empirical probability.

Whether in inter-sample comparisons at the same moment of the process (comparison between the experimental and control groups at the beginning of the research), or whether we perform the comparative study of the results of the same sample at two different moments of the study, before and after our experimental manipulation, in both situations the phenomenon may occur that even when N is equal in the two groups — experimental and control — and the two measurements, before and after the experimental manipulation, the number of positions in the ranking, N^o , might not be equal, regardless of the fact that the magnitude N remains the same.

At the moment when the total number of positions N^o in the ranking must be compared — that is, the sample N^o of the ranking — between two samples with different sample N^o of the ranking, regardless of whether the number of subjects or options remains constant, the comparison of samples N^o of the ranking between two samples automatically requires the consideration of the *effect of N^o* over the data, insofar as, just as in the Second Method, in the same way that the magnitude of the sample of subjects or options N has direct consequences on the estimation of empirical probabilities, likewise differences in the sample N^o of the ranking will have direct consequences for data interpretation. Insofar as, as the sample of positions in the ranking decreases, the expectations of reliability in the results increase.

Just as N has an effect on the sample, the sample N^o will also have effects on the sample, in the sense that when we later move on to the critical intermeasurement study of a subject or option, the positions occupied by the same subject or option in two different samples or measurements of different N^o may only be compared if the Effect of N^o on the position occupied by that subject or option in the second measurement is taken into consideration.

In the particular case of ranking studies, the Effect of N^0 will be equal to the quotient of the number of positions N^0 in the ranking of the first measurement divided by the number of positions N^0 in the ranking of the second measurement, and the result is multiplied by the position corresponding to the subject or option in the second measurement that one wishes to compare relative to its position in the first measurement.

Effect of $N^0 = N^0/1^a : N^0/2^a$

Effect of N^0 on a ranking position in the second measurement

$$n^0/2^a \cdot (N^0/1^a : N^0/2^a)$$

$n^0/2^a$ = position in the second measurement

$n^0/1^a$ = position in the first measurement

$N^0/1^a$ = total number of positions in the ranking in the first measurement

$N^0/2^a$ = total number of positions in the ranking in the second measurement

It can also be applied to interindividual studies. Let us suppose that the same medical treatment is applied to different samples of subjects over a period of time, regardless of whether the passage of time is seconds, minutes, days, months, years, or even a century or more; if both samples of patients have been subjected to the same medical treatment, it would be important evidence to compare our results, and our position in the ranking, and determine which patient occupies the N^0 position in the ranking.

64. Inter-measurement Studies in Ranking Studies When the Last Position Is the Ideal

In experimental models, it is rare that the expected results are achieved in a first measurement; in fact, what is normal in our work is always to take a first measurement before beginning the experimental phase in order to know the initial position of our subjects or options, and then, once the experimental variable begins to be applied, to carry out successive measurements to determine the evolution of the subjects or options as a function of the real experimental manipulation.

Thus, in essence, the development of inter-measurement studies will be necessary from the very moment a collection of measurements begins to be generated.

At the moment collections of data are obtained, the comparison of data between different measurements is considered an inter-measurement study, and in this type of study the Effect of N^0 will be crucial in order to compare concrete ranking positions between two different N^0 samples.

If in a first measurement the magnitude N^0 is greater than the magnitude N^0 of the second measurement, and one wishes to compare the particular position n^0 of a subject or option in the second measurement relative to the position n^0 of that same subject or option in the first measurement, then the Effect of N^0 must be applied to the particular position n^0 of that subject or option in the second measurement so that it can be compared with the position n^0 of that subject or option in the first measurement.

Insofar as, both in the ascending sense of the ranking, where the maximum value occupies the last position, and in the descending sense of the ranking, where the last position occupies the minimum, if the object of study is the ascending study from lower to higher, or descending from higher to lower, whether the object of study is negative bias — where the empirical probability tends to the minimum (last position in the descending ranking) — or positive bias — where it tends to the maximum (last position in the ascending ranking) — in both models, insofar as the position n^0 of a subject or option occupies a superior position, it will signify a greater tendency of the subject or option toward its ideal, regardless of whether its ideal is the minimum or the maximum.

In both cases, positive or negative bias, the inter-measurement critique will aim, by applying the Effect of N^0 to the second measurement, to study how, in the second measurement, the position of the subject or option in the ranking, after applying the Effect of N^0 , is superior to the position of that subject or option in the first measurement.

The inter-measurement studies presented are, first, the *Ranking Level*, where positions are directly compared taking into account the Effect of N^0 in the second measurement; and *Ranking Significance*, the adaptation of Ranking Level to a critical rational contrast test.

Ranking Level in ascending growth or descending decrease

$$[n^o/2^a \cdot (N^o/1^a : N^o/2^a)] - n^o/1^a = \text{positive is accepted}$$

$n^o/2^a$ = position of the second measurement

$n^o/1^a$ = position of the first measurement

$N^o/1^a$ = total number of positions in the ranking in the first measurement

$N^o/2^a$ = total number of positions in the ranking in the second measurement

Ranking Significance in ascending growth or descending decrease

$$\{ [n^o/2^a \cdot (N^o/1^a : N^o/2^a)] - n^o/1^a \} - nc = \text{zero or positive is accepted}$$

$$nc = (N^o/1^a - n^o/1^a) \cdot (X : 100)$$

X = percentage of reliability

65. Inter-measurement studies in ranking studies when the ideal is to be the first position

In *Impact of the Defect*, the categories are ordered from lesser defect to greater defect in ascending order; therefore, the last category corresponds to the maximum possible defect. In *Effective Distribution*, the categories of efficacy and efficiency are ordered in ascending order, so that maximum efficacy and efficiency correspond to the last position of the Effective Distribution.

The logic behind ascending ordering of defects or efficacy is simple: since the last position always occupies the maximum position, the maximum GP, Weighted Gravity, corresponds to the GP of the last position; the last position occupies the maximum GP

$$GP = ((n^o = N^o) : N^o) = 1.$$

If the maximum GP is when $GP = 1$, the same occurs with IP: the maximum Weighted Importance, IP, is when IP is equal to one, $IP = ((n^o = N^o) : N^o) = 1$. Thus, the aim of Effective Distribution (EI) and Impact of the Defect (ID) is to order the positions in ranking so that, in the end, the maximum possible weightings are $GP = 1$, $IP = 1$, in ID and EI.

However, situations may arise in ranking studies such as those analysed in section 18 of *Introduction to Impossible Probability*, in which the ordering of the factors in the ranking may follow different models, not necessarily placing the maximum values in the last category.

If such a case were to occur, then the ordering would be different. If it happens that the categories are ordered in such a way that the first category corresponds to the first position, as in the Olympics, or the first category corresponds to the minimum, if the first position becomes occupied by the maximum or minimum depending on the ideal — in any case, the ideal is to be the first position — then the ordering of the factors in Ranking Level and Significance should be adapted to fit the new order of categories in the ranking, with the first position becoming the ideal.

Ranking Level when the maximum in growth or minimum in decrease is the first position

$$n^0/1^a - [n^0/2^a \cdot (N^0/1^a : N^0/2^a)] = \text{positive is accepted}$$

Ranking Significance when the maximum in growth or minimum in decrease is the first position

$$\{ n^0/1^a - [n^0/2^a \cdot (N^0/1^a : N^0/2^a)] \} - nc = \text{zero or positive is accepted}$$

$$nc = n^0/1^a \cdot (X : 100)$$

X = percentage of reliability

“In the fell clutch of circumstance
 I have not winced nor cried aloud.”

“En las garras crueles de las circunstancias
 no he pestañado ni he llorado en voz alta.”

Invictus

William Ernest Henley

Index

Introduction to *Introducción a la Probabilidad Imposible*: Probability Statistics, or Statistical Probability

Analysis of *Introducción a la Probabilidad Imposible*

Presentation of the English Version of *Analysis of Introduction to Impossible Probability*

1. Statistics of probability
2. Empirical probability
3. Differences between Second Method and traditional statistics
4. Types of universe and types of sample
5. Infinite universes and theoretical dispersion
6. Infinite universes and sample representativeness error
7. The importance of N
8. Empirical individual statistics
9. Relative statistics
10. Theoretical individual statistics
11. Empirical sample statistics
12. Theoretical sample statistics
13. Normal Bias Level and object of study
14. Omega models
15. Critical probability and inferential statistics
16. Critical Bias Level, Validity and individual Signification
17. Critical Sample Level and Sample Signification
18. Rational criticism individual and sample in omega models
19. Theoretical individual critical proportions
20. Empirical individual critical proportions
21. Empirical individual critical proportions

22. Empirical sample critical proportions
23. Probability of Equality and Probability of Bias, Positive or Negative
24. Rational criticism of the typical score
25. Critical Proportions of normal Typical Score
26. Rational Criticism of Typical Score under omega conditions
27. Ideal Similarity Critical of Typical Score
28. Criticism of magnitude N and the effect of N
29. Necessary variation and critical possible variation
30. Types of empirical error and empirical reliability
31. Necessary variation and empirical possible variation
32. Individual projection
33. Sample projection
34. Individual and sample prognosis
35. Intermeditational individual studies
36. Reliable Signification in intermeditational individual study
37. Critical Intermeditational individual proportion
38. Intermeditational Intraindividual Descriptive Statistics
39. Intermeditational Sample Criticism
40. Omega intermeditational individual studies on a critical ratio
41. Omega intermeditational sample studies on a critical ratio
42. Third modality of empirical error and empirical reliability in omega models
43. Omega intermeditational individual studies on empirical reliability
44. Omega intermeditational sample studies on empirical reliability
45. Rational intermeditational individual omega criticism
46. Rational intermeditational sample omega criticism
47. Omega Critical Sample Proportion
48. Individual Defect Impact
49. Theoretical individual statistics in Defect Impact
50. Criticism of the Individual Defect Impact
51. Descriptive statistics of Defect Impact
52. Theoretical sample statistic of Defect Impact, $1/N^0$

53. Sample criticism of Defect Impact
54. Rational Intermeditational criticism individual and sample
55. Effective Distribution: Individual Effectiveness
56. Descriptive statistics of Individual Effectiveness
57. Theoretical Statistics of Individual Effectiveness
58. Rational Intermeditational Criticism of Individual Effectiveness
59. Sample criticism in Effective Distribution
60. The statistical ranking of empirical probabilities
61. Theoretical statistics of homogeneity in ranking
62. Rational criticism of ranking homogeneity
63. N° Effect of ranking
64. Intermeditational studies in ranking studies when the last position is the ideal
65. Intermeditational studies in ranking studies when the ideal is to be the first position

“And it is in the light of dawn
where your soul is reflected,
hidden among the rye
and watched over by the guardian’s tower.
Here lies the key, if you know how to use it.”

Military Project
Rubén García Pedraza
6 December 2025, The Mitchell Library, Glasgow, Scotland
12:31

Translation finished at 16:35,
9 December 2025, Belfast, Ireland.
Imposiblenever@gmail.com